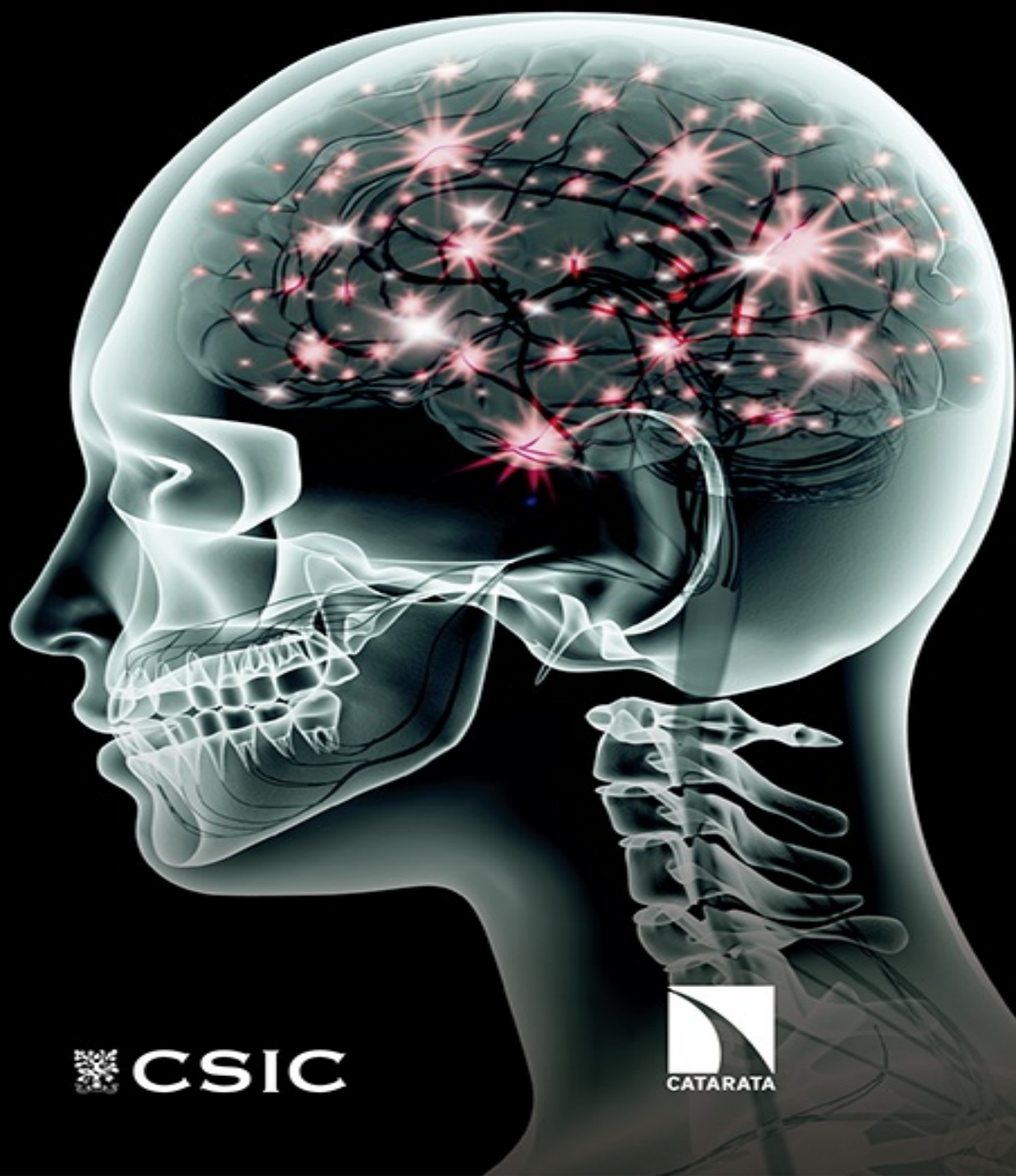


¿QUÉ SABEMOS DE?

# Neuromatemáticas

## El lenguaje eléctrico del cerebro

José María Almira y Moisés Aguilar-Domingo



 **CSIC**

  
CATARATA

¿QUÉ SABEMOS DE?

# Neuromatemáticas

## El lenguaje eléctrico del cerebro

José María Almira y Moisés Aguilar-Domingo



 **CSIC**

 **CATARATA**

# **Neuromatemáticas**

El lenguaje eléctrico del cerebro

José María Almira y Moisés Aguilar-Domingo



Colección ¿Qué sabemos de?

COMITÉ EDITORIAL

CONSEJO ASESOR

PILAR TIGERAS SÁNCHEZ, DIRECTORA

BEATRIZ HERNÁNDEZ ARCEDIANO, SECRETARIA

RAMÓN RODRÍGUEZ MARTÍNEZ

JOSE MANUEL PRIETO BERNABÉ

ARANTZA CHIVITE VÁZQUEZ

JAVIER SENÉN GARCÍA

CARMEN VIAMONTE TORTAJADA

MANUEL DE LEÓN RODRÍGUEZ

ISABEL VARELA NIETO

ALBERTO CASAS GONZÁLEZ

JOSÉ RAMÓN URQUJO GOITIA

AVELINO CORMA CANÓS

GINÉS MORATA PÉREZ

LUIS CALVO CALVO

MIGUEL FERRER BAENA

EDUARDO PARDO DE GUEVARA Y VALDÉS

VÍCTOR MANUEL ORERA CLEMENTE

PILAR LÓPEZ SANCHO

PILAR GOYA LAZA

ELENA CASTRO MARTÍNEZ

ROSINA LÓPEZ-ALONSO FANDIÑO

MARÍA VICTORIA MORENO ARRIBAS

DAVID MARTÍN DE DIEGO

SUSANA MARCOS CELESTINO

CARLOS PEDRÓS ALIÓ

MATILDE BARÓN AYALA

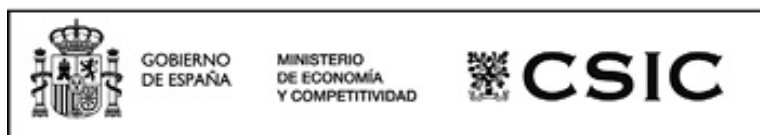
PILAR HERRERO FERNÁNDEZ

MIGUEL ÁNGEL PUIG-SAMPER MULERO

JAIME PÉREZ DEL VAL

CATÁLOGO GENERAL DE PUBLICACIONES OFICIALES

[HTTP://PUBLICACIONESOFICIALES.BOE.ES](http://publicacionesoficiales.boe.es)



Diseño gráfico de cubierta: Carlos Del Giudice

Fotografía de cubierta: © iStock/Thinkstock

© José María Almira y Moisés Aguilar-Domingo, 2016

© CSIC, 2016

© Los Libros de la Catarata, 2016

Fuencarral, 70

28004 Madrid

Tel. 91 532 20 77

Fax. 91 532 43 34

[www.catarata.org](http://www.catarata.org)

ISBN (CSIC): 978-84-00-10115-2

ISBN ELECTRÓNICO (CSIC): 978-84-00-10116-9

ISBN (CATARATA): 978-84-9097-219-9

ISBN ELECTRÓNICO (CATARATA): 978-84-9097-220-5

NIPO: 723-16-254-5

NIPO ELECTRÓNICO: 723-16-255-0

DEPÓSITO LEGAL: M-33.566-2016

IBIC: PDZ/PSAN

ESTE LIBRO HA SIDO EDITADO PARA SER DISTRIBUIDO. LA INTENCIÓN DE LOS EDITORES ES QUE SEA UTILIZADO LO MÁS AMPLIAMENTE POSIBLE, QUE SEAN ADQUIRIDOS ORIGINALES PARA PERMITIR LA EDICIÓN DE OTROS NUEVOS Y QUE, DE REPRODUCIR PARTES, SE HAGA CONSTAR EL TÍTULO Y LA AUTORÍA.

## CAPÍTULO 1

# Electroencefalograma y potenciales evocados: el lenguaje eléctrico del cerebro

### Qué es un electroencefalograma, cómo se graba y para qué sirve

El cerebro es, sin duda, el sistema más complejo al que podemos enfrentarnos. También es probablemente aquel cuyo conocimiento supone el reto más importante que podamos afrontar. Somos nuestro cerebro. En tiempos remotos esto no estaba claro y numerosas teorías colocaban el centro de nuestro ser, el alma, la conciencia, el pensamiento, en otras partes del cuerpo, como el corazón, el hígado o los intestinos. Aristóteles, padre de la lógica, pensaba que el cerebro no participa para nada en nuestro “pensamiento”. Se convenció de ello tras realizar algunos experimentos con animales, comprobando que este órgano es por sí mismo insensible (realizó cortes sin que observara respuesta dolorosa por parte de los animales sometidos a dicha experiencia). Como consecuencia de esto, propuso que el corazón es el órgano de las sensaciones. Posteriormente, Hipócrates, en sus estudios sobre la epilepsia (enfermedad a la que se adjudicaba un carácter sagrado), se basó en las observaciones anatómicas de Acmeón y sus discípulos, que habían establecido la existencia de nervios que, con continuidad, parten del cerebro y alcanzan todas las partes del cuerpo, para concluir que este mal, que provoca convulsiones en diferentes partes del cuerpo, tiene en realidad su origen en un mal funcionamiento del cerebro. Además, afirmó que los ojos, manos, pies, lengua y oídos se comportan de acuerdo a ciertas instrucciones que reciben del cerebro, y que este es el centro de nuestro entendimiento. Sin

embargo, hubo que esperar a Claudio Galeno para que se aportaran suficientes pruebas de que, en efecto, el cerebro y la médula espinal (es decir, lo que actualmente denominamos sistema nervioso central) tienen un papel fundamental en nuestra conducta.

Evidentemente, existen numerosos puntos de vista que nos permiten acercarnos a una mejor comprensión del cerebro. Podemos, por ejemplo, estudiar los procesos bioquímicos, así como los procesos bioeléctricos que tienen lugar en dicho órgano y que son consustanciales a su correcto funcionamiento. Podemos centrarnos en aspectos fisiológicos y/o anatómicos y relacionar las distintas regiones que componen el encéfalo con diversas funciones neurológicas. Podemos centrar nuestra atención en las unidades mínimas que componen el cerebro: las neuronas y las células gliales, y estudiar su comportamiento y su estructura. Podemos, por supuesto, plantearnos el estudio de las redes neuronales, intentando explicar cómo se activan y qué efecto tienen para generar las distintas actividades que se adjudican hoy en día al cerebro. Podemos abstraernos de todo lo anterior e intentar explicar, desde un punto de vista más teórico, los aspectos básicos del cerebro como máquina, etc. Ninguno de estos enfoques, por sí mismo, bastará para alcanzar un entendimiento completo del cerebro. Cualquiera de ellos basta para llenar una vida dedicada por completo al estudio.

En este breve texto introducimos algunos aspectos básicos relacionados con el lenguaje eléctrico del cerebro, tal como se percibe desde el exterior del cuero cabelludo. Dicha actividad eléctrica, que se produce con amplitud de microvoltios, se recoge mediante el uso de electrodos repartidos homogéneamente en la cabeza, dando lugar a los llamados electroencefalogramas (EEG). Es decir, un EEG está formado por un conjunto relativamente pequeño de señales eléctricas, medidas en microvoltios, que han sido grabadas por diferentes sensores que se colocan en el exterior del cuero cabelludo. Se pueden grabar electroencefalogramas con sensores que tengan contacto directo con la piel o incluso con el encéfalo, tanto a nivel superficial

como profundo. Sin embargo, en este texto nos hemos centrado exclusivamente en el caso en el que la grabación se realiza desde el cuero cabelludo. El número de electrodos utilizados puede variar, desde la decena hasta el orden de trescientos.

Los primeros EEG humanos fueron grabados en Jena, Alemania, por H. Berger en la década de 1920. Para tomar sus medidas, utilizó un galvanómetro de cuerda, que es un instrumento poco preciso cuando las intensidades de corriente eléctrica son muy pequeñas. Afortunadamente, pronto surgió el tubo de rayos catódicos, que permite realizar ampliaciones importantes, sin dar lugar a distorsiones y, por tanto, es un instrumento mucho más apropiado para medir electroencefalogramas. Evidentemente, el paso del tiempo ha ido aportando numerosas mejoras tecnológicas y en la actualidad existen electrodos que realizan su función con enorme precisión y eficiencia. En particular, es necesario destacar aquí la elevada resolución temporal que tienen estos aparatos.

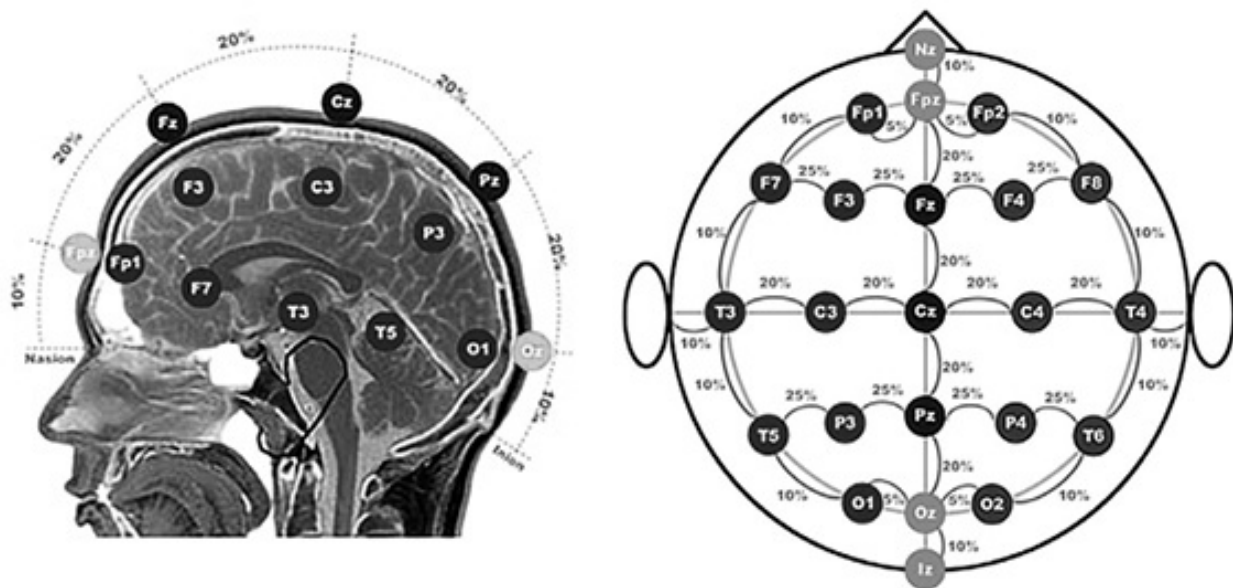
La actividad recogida por un electrodo cambia de forma sensible cuando modificamos su posición. Esto es natural, pues aunque el electrodo recibe información que proviene en realidad de numerosas partes del encéfalo, la actividad que se genere en la región inmediatamente más próxima al mismo tendrá un efecto predominante. Si estamos en medio de una fiesta con mucha gente hablando, podremos escuchar mucho mejor a aquellos interlocutores que se encuentren cerca de nosotros que al resto, a los cuales seguiremos oyendo, pero lo que dicen nos resultará indiscernible o incluso lo interpretaremos como ruido. Algo similar sucede con nuestro electrodo, pero él solo “escucha” potenciales eléctricos. Es por ello que, para que un EEG pueda recoger suficiente información, útil para posibles usos clínicos, es absolutamente necesario utilizar un número suficientemente elevado de electrodos en nuestra grabación. Desafortunadamente, un número excesivo de electrodos puede también resultar perjudicial, pues los datos recogidos por cada electrodo afectan a un área no especialmente pequeña, por lo que la resolución espacial del EEG no puede superar cierto umbral. Además, existe la posibilidad de interferencias.



Si queremos utilizar la información recogida por un EEG para realizar cualquier tipo de diagnóstico es absolutamente necesario que hayamos generado previamente un patrón de “normalidad” en humanos “sanos”, contra el cual debemos contrastar los datos obtenidos. Para poder hacer esto, se impone crear una base de datos de pacientes sanos (así como de pacientes con distintas patologías) y esta debe realizarse de forma homogénea. Es decir, debemos disponer del mismo número de electrodos, colocados en idéntica posición, y con idéntico montaje (aunque los electrodos se coloquen igual, hay que definir cierto sistema de referencia que fija la manera en la que tomamos las medidas, pues no medimos el potencial eléctrico tal cual, sino su diferencia respecto de un cierto nodo de referencia que hay que fijar) así como definir un protocolo específico para la toma de datos. Existen numerosos protocolos para la medición de EEG, pero hay uno que destaca sobre los demás debido a que es el más extendido, aquel con el cual se han realizado más bases de datos y el que mejor contrastado está. Es el sistema internacional 10-20 y es el que nosotros usaremos a lo largo de este libro. Se basa en la relación existente entre la localización de un electrodo y la región del cerebro que hay bajo la misma. Los números 10 y 20 hacen referencia a la forma en la que se toman las distancias entre electrodos adyacentes, que son o el 10% o el 20% de la distancia total de delante a atrás (es decir, entre los puntos nasión e inión) o desde la derecha a la izquierda (es decir, entre los puntos preauriculares), en el cráneo.

**FIGURA 1**  
Sistema internacional 10-20 para el registro de un EEG.





Para colocar los electrodos caben dos posibilidades: bien realizamos las medidas correspondientes en el sujeto y procedemos a colocar cada electrodo de forma individual, ayudados de varias cintas que nos faciliten su fijación, o bien utilizamos un gorro flexible, que tiene un orificio en cada uno de los nodos pertinentes, ya medido, donde colocamos cada uno de los electrodos. En ambos casos es necesario que se aplique un cierto gel conductor que permita un contacto completo entre cabeza y electrodo. Normalmente, el uso de gorros reduce el tiempo necesario para la colocación de los electrodos y, además, resulta más cómodo para los pacientes. Es muy importante que este trabajo se realice a conciencia, pues una medición imprecisa o artefactada nos dará problemas para el diagnóstico clínico objetivo. Si no tomamos bien los datos del EEG, todo el trabajo posterior no servirá de nada.

Aún habiendo fijado los electrodos en sus respectivas posiciones, existen diferentes formas posibles de grabación del EEG (o montajes). Esto es porque en cada uno de los puntos donde colocamos un electrodo es conveniente grabar la diferencia de potencial eléctrico entre dicho punto o nodo y otro nodo de referencia que debemos fijar de antemano y, por supuesto, existen numerosas formas de elegir estos nodos de referencia. Así, en el sistema 10-20 no existe un

único montaje posible, sino varios, y hay que especificar cuál de ellos se está usando. Por otra parte, la colocación de los electrodos no es arbitraria, sino que tiene bases fisiológicas, pues desde hace tiempo se conocen las principales funciones neurocognitivas gestionadas por las distintas regiones del encéfalo y, de hecho, existe un atlas del cerebro bastante completo en el que dichas funciones están especificadas. Además, la excelente resolución temporal con la que se graba el EEG implica que este recoge la actividad eléctrica del encéfalo con total precisión en las escalas temporales a las que funciona, que son las mismas escalas con las que el sujeto debe interaccionar con el medio. Así pues, las medidas del EEG tienen un valor especial cuando queremos utilizarlas para interpretar los diversos procesos neurocognitivos que existen. Se sigue que la actividad en las distintas regiones del encéfalo, tal como es captada por los electrodos del EEG, es susceptible de interpretación clínica.

### Actividad bioeléctrica de la neurona

Las neuronas, ¿cómo generan electricidad? La respuesta es que se comportan, en gran medida, como una pila electroquímica. Recordemos que una pila posee dos celdas separadas, las cuales contienen una solución de electrolitos (iones), cada una a diferente concentración. Cuando las celdas se comunican a través del contacto de ambas con un mismo material conductor, aparece un flujo de corriente eléctrica, que va desde la celda con mayor concentración de carga hacia la celda con menor concentración. Si asumimos que en ambas celdas solo está presente un tipo de ión y denotamos por  $V$  la diferencia de potencial eléctrico

$$V = \frac{RT}{zF} \ln \frac{C_1}{C_2}$$

generada, se sabe (ley de Nernst) que esta satisface la ecuación , donde  $C_1, C_2$  representan las concentraciones iónicas de la solución a uno y otro lado de la membrana,  $z$  es la valencia del ión bajo consideración,  $F$  es la constante de Faraday,  $R$  es la constante de los gases y  $T$  es la temperatura

absoluta de la solución. Finalmente,  $\ln(t)$  representa el logaritmo neperiano. Una vez esto ha quedado establecido, la observación clave consiste en constatar el hecho de que las células nerviosas mantienen, con el medio que las circunda, una diferencia de concentración de ciertas sales ionizadas y por tanto, estas son susceptibles de comportarse como una pila electroquímica. Es más, una parte fundamental de la energía consumida por estas células se emplea en mantener esta diferencia de concentración por encima de un determinado umbral, cosa que deja de suceder solo cuando la célula muere o cuando se interrumpe su metabolismo.

La membrana de la célula, que es permeable en distintos grados a las diferentes sales que se encuentran en el citoplasma así como en el exterior de la célula, se ocupa de mantener aproximadamente constante la diferencia de potencial entre ambas partes, cuando esta no está excitada eléctricamente. Para lograr esto, utiliza unos canales de paso que permiten, de manera selectiva, que ciertos iones viajen de un lado a otro de la célula, en conjunción con una serie de reacciones químicas que se producen gracias a la presencia de cierta proteína transmembranal (que usa energía del metabolismo para cambiar las concentraciones de sodio y potasio) y que reciben el nombre de “bomba de sodio-potasio-ATPasa”, pero no queremos entrar aquí en excesivos detalles. Dicha diferencia de potencial recibe el nombre de potencial de reposo o potencial de membrana. El potencial que existe en el interior de la célula, si no está eléctricamente estimulada, es negativo respecto del existente en el exterior. Este potencial se mantiene mientras los iones no puedan salir de la célula ni entrar a ella. Si entran o salen, su concentración relativa varía y la magnitud del potencial de reposo se altera, dando lugar a un potencial de acción o impulso eléctrico.

Así pues, cuando la neurona se excita eléctricamente, lo que sucede es que se produce una variación brusca del potencial de reposo, que da lugar a una onda de descarga eléctrica que, mientras se va realizando, ocupa en la membrana de la célula solo una zona restringida, y luego se propaga por la superficie que aún

queda libre en el axón. Es importante, además, observar que la velocidad de propagación de esta onda depende del diámetro del axón (a mayor diámetro, mayor velocidad). Estas sacudidas, excitaciones o potenciales de acción son eventos que están localizados tanto espacial como temporalmente, lo cual los convierte en un mecanismo muy útil para transmitir información entre distintos puntos. Al producir excitaciones en momentos y lugares apropiados, desde donde se propagan, las neuronas se envían mensajes entre sí o los envían a otros tejidos corporales, como músculos (regulando de este modo su contracción) o glándulas (cuya función secretora también se controla así). Esto es esencialmente lo que sucede con una única neurona.

Evidentemente, el potencial eléctrico recogido en el EEG no proviene de la actividad eléctrica de células aisladas sino que, por el contrario, requiere que miles de células nerviosas se activen simultáneamente, de manera conjunta y sincronizada, para generar un potencial eléctrico lo suficientemente elevado como para poder captarlo desde el exterior del cuero cabelludo. Y aun así, en ningún caso podremos recoger en los electrodos la totalidad de la actividad eléctrica del encéfalo. Por ejemplo, hay enormes posibilidades de que diferentes potenciales, al encontrarse, se anulen entre sí. También es posible que la dirección de propagación de la corriente eléctrica, si sigue una trayectoria oblicua respecto del cráneo, no llegue al exterior del mismo con la suficiente intensidad como para que esta pueda ser grabada. Así, se asume que los EEG recogen la actividad eléctrica que se propaga perpendicularmente (o casi perpendicularmente) al cráneo. A pesar de esta importante limitación, resulta que dicha actividad eléctrica ha demostrado contener información sustancial para la comprensión de numerosos aspectos relacionados con nuestros procesos neurocognitivos y, de hecho, puede ser utilizada para el diagnóstico (y el tratamiento) de diversos desórdenes o trastornos mentales.

## Los ritmos del cerebro

Existen muchos tipos de neuronas. De hecho, hay neuronas con muy diversos

tamaños, diferentes morfologías, composición química, etc. Algunas se pasan la vida en permanente actividad, comunicándose sin cesar con sus compañeras. Otras, sin embargo, permanecen largo tiempo en silencio y, de pronto, se excitan. Algunas se excitan con rápidas ráfagas de impulsos eléctricos, y otras son más lentas. Sin embargo, existe un concepto, una categoría, que permite dividir todos estos tipos de neuronas esencialmente en dos clases —y cada neurona pertenece a una de estas clases—; o bien nuestra neurona tiene la capacidad de excitar a aquellas neuronas con las que está conectada, o bien tiene la capacidad de inhibirlas. Concretamente, las neuronas del tipo “excitador” tienen la capacidad de despolarizar las neuronas postsinápticas a las que están conectadas, lo cual hace que el umbral necesario para activar sus potenciales de acción disminuya y, por tanto, sean más fáciles de excitar. Por otra parte, las neuronas de tipo “inhibidor” hiperpolarizan a las neuronas postsinápticas a las que están conectadas, provocando que estas resulten más difíciles de ser excitadas.

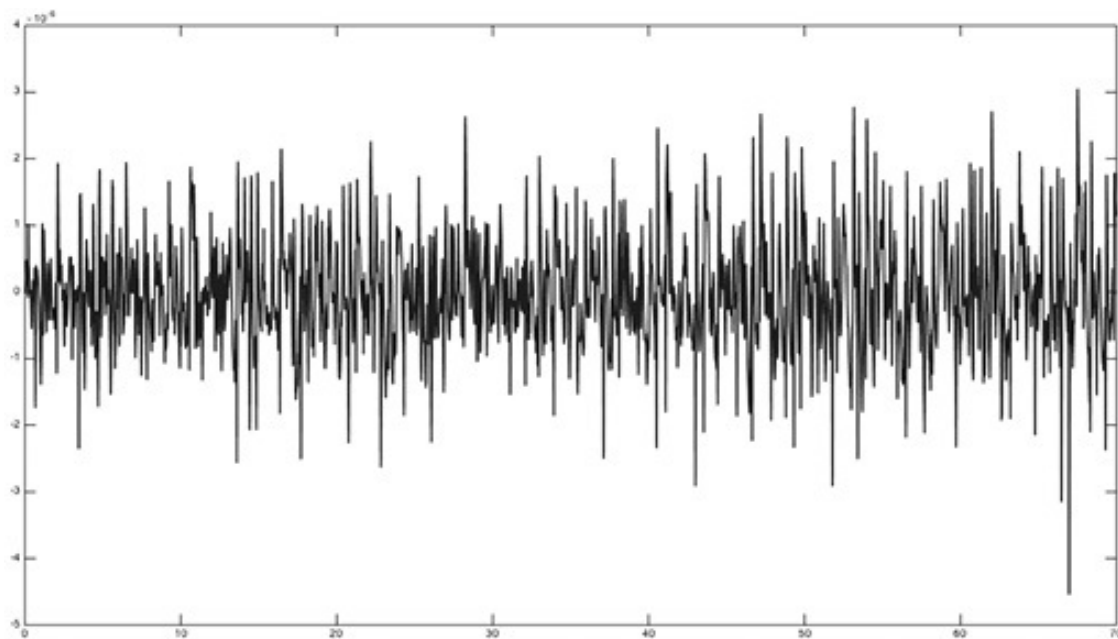
Ahora bien: todas las neuronas están densamente interconectadas, y lo están tanto con neuronas excitadoras como inhibidoras. Supongamos ahora, por ejemplo, que abrimos los ojos y la luz penetra a través de ellos, con numerosos matices de color, estimulando fuertemente nuestro córtex visual. Entonces numerosas neuronas excitadoras activarán a todas aquellas neuronas con las que están conectadas y, entre estas, habrá muchas neuronas que son también excitadoras. Esto incrementará el potencial eléctrico generado en esta zona del córtex. Ahora bien, nuestras neuronas están también conectadas a neuronas inhibidoras. Ellas, al activarse, comenzarán su trabajo de inhibición de las neuronas conectadas postsinápticamente a ellas. De este modo llegará un momento en el que el potencial eléctrico dejará de crecer y, por el contrario, comenzará a decrecer. Las neuronas inhibidoras también serán inhibidas y, con el paso del tiempo, llegará un momento en el que nuevamente la actividad de las neuronas excitadoras tome ventaja sobre las neuronas inhibidoras. Lo que aquí se produce es, por tanto, una permanente lucha entre activación e inhibición, una lucha que causará, en definitiva, que el potencial eléctrico oscile de arriba a

abajo, formando ciclos. Este es el mecanismo principal que genera oscilaciones neuronales. Por supuesto, no es el único mecanismo neuronal que tiene la capacidad de producir oscilaciones o ciclos. Es posible lograr comportamientos oscilatorios mediante el uso de redes neuronales que son solo excitadoras, por ejemplo. De hecho, las oscilaciones se producen con bastante facilidad, incluso en redes neuronales bastante pequeñas. Lo raro sería no encontrar comportamientos oscilatorios en la medida del potencial eléctrico generado por una red neuronal.

Las oscilaciones que podemos observar en el potencial registrado por un EEG varían desde un poco menos de 0,5 Hz hasta por encima de los 400 Hz. Sin embargo, no encontramos presencia de todas las frecuencias posibles en este intervalo, ni tampoco son igualmente comunes unas frecuencias y otras.

FIGURA 2

Señal registrada durante unos pocos segundos por un electrodo del EEG.



Con los métodos rudimentarios de los que disponía, Berger detectó que, cuando se registra la actividad eléctrica en la región occipital a un paciente que, estando despierto, cierra los ojos, “el electroencefalograma representa una curva

continua con oscilaciones continuas de las que podemos distinguir en primer lugar un grupo de ondas de mayor amplitud cuya duración es de 90 milisegundos, y un segundo grupo de ondas, más pequeñas, cuya duración media es de 35 milisegundos. Las ondas mayores miden entre 150 y 200 microvoltios”.

En particular, el campo eléctrico generado por la descarga de millones de neuronas en el córtex cerebral es unas 10.000 veces inferior al generado por una pila de 1,5V. Berger denominó a las ondas del primer tipo (que oscilan a 10 Hz, es decir, 10 veces por segundo), ondas  $\alpha$ , y a las del segundo tipo, las llamó ondas  $\beta$ .

Actualmente se distinguen esencialmente cinco grupos de ondas en el EEG, atendiendo a su comportamiento en frecuencias. Estos grupos se han bautizado con letras griegas. Son las ondas  $\delta$  (entre 0,5 y 2 Hz),  $\theta$  (en torno a los 6 Hz),  $\alpha$  (cerca de los 10 Hz),  $\beta$  (cerca de los 30 Hz) y, por último,  $\gamma$  (en torno a 50 Hz). Con toda seguridad existen importantes razones, de carácter biológico, fisiológico, etc., por las que aparecen oscilaciones precisamente en estos rangos de frecuencias. Por ejemplo, hay neuronas cuyos canales de conductividad — aquellos que permiten el paso de iones de sodio, potasio, etc., del interior al exterior de la célula, y viceversa—, se activan aproximadamente cada 150 milisegundos. Si estas neuronas son estimuladas, dichos canales se activarán con mayor frecuencia, alrededor de cada 100 milisegundos, pero nunca a una velocidad mayor. Así, los potenciales de acción generados por estas neuronas aparecerán con una frecuencia que ronda los 10 Hz y esto causará la aparición de ondas  $\alpha$  en el correspondiente EEG. Otros tipos de neuronas podrán ser la causa de que se registren ondas en otros anchos de banda en el EEG.

Resulta muy significativo que, aunque los encéfalos de las distintas especies han evolucionado considerablemente a nivel fisiológico, cuando medimos el EEG de especies evolutivamente alejadas, como un ratón y un humano, y lo vemos en términos de frecuencias, apenas hay diferencias significativas, pues ambos cerebros generan las mismas frecuencias. ¿Debería sorprendernos esto? ¡No! Lo cierto es que ambas especies comparten el mismo entorno y, por tanto,



deben ser capaces de responder eficientemente ante sucesos que se producen en la misma escala temporal, por lo que sus sistemas nerviosos necesitan procesar y responder a estímulos que les llegan con las mismas frecuencias ¡y eso es lo que hacen!

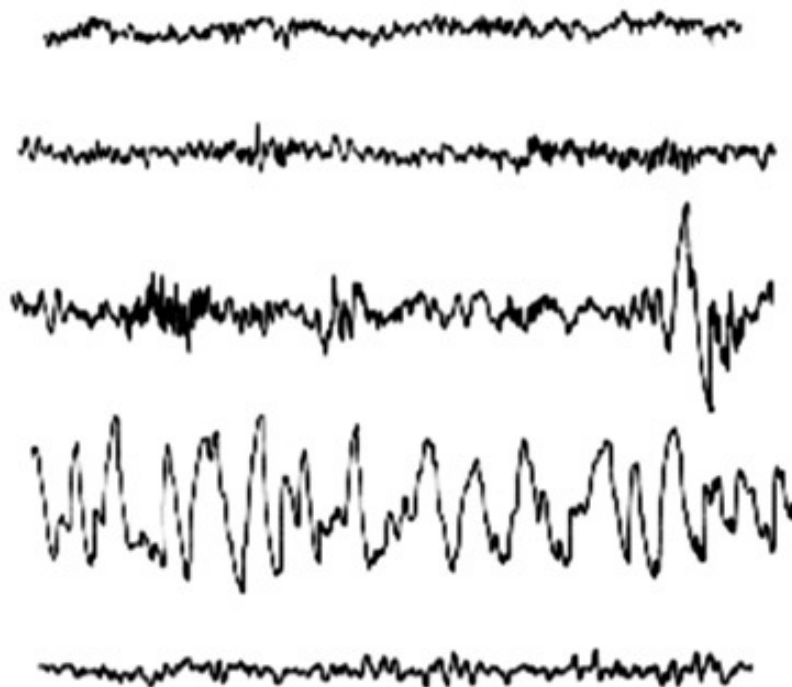
Que las distintas frecuencias que aparecen en un EEG son significativas desde un punto de vista neurocognitivo es algo que podemos deducir de numerosas fuentes. Por ejemplo, se sabe que las ondas  $\alpha$  son importantes en la actividad eléctrica espontánea registrada por un EEG en el córtex visual humano. Si sometemos a estímulos dicha zona del cerebro, por ejemplo, mediante la presentación de luz estroboscópica, resulta que estos estímulos viajan desde la retina, a través del tálamo, hasta el córtex visual. Si los estímulos llegan allí precisamente en el momento en que las oscilaciones espontáneas en la banda  $\alpha$  se encuentran en una fase asociada a una excitación elevada (es decir, cerca de un pico de actividad espontánea, un máximo), entonces las neuronas que reciben dicho estímulo podrán responder al mismo con gran facilidad, activando sus potenciales de acción, lo cual aumentará la amplitud del EEG registrado en ese ancho de banda en ese instante. Si por el contrario los estímulos llegan en un momento (o fase) en el que las neuronas que los reciben están más inhibidas (es decir, cerca de un valle, un mínimo), entonces dichos estímulos tendrán mayor dificultad a la hora de activar potenciales de acción en dichas neuronas y, por tanto, serán difícilmente percibidos, o percibidos solo de forma muy débil. Este principio se ha utilizado con enorme éxito para “entrenar” al cerebro con técnicas de estimulación visual, basándonos en la idea de que someter al cerebro a un estímulo de forma intermitente y repetida logra que la zona del cerebro encargada de procesar la información recibida (en nuestro caso, el córtex visual) comience a oscilar a la misma frecuencia con la que se está realizando la estimulación. De hecho, numerosas experiencias corroboran la existencia de importantes ligaduras entre la percepción y las oscilaciones registradas por los EEG. Lo mismo sucede con otros procesos neurocognitivos, como la atención, la impulsividad, o la acción. Por ejemplo, se sabe que un incremento de la

actividad  $\alpha$  (en términos de amplitud) está relacionado con la desconexión de ciertas regiones del cerebro (normalmente, del córtex somatosensorial) cuando estas no son necesarias o cuando su actividad puede generar una distracción que no se desea. Es decir, un incremento de la actividad  $\alpha$  en ciertas zonas del córtex está vinculado a un proceso de inhibición activa. Para lograr una mayor concentración, podría ser necesario desactivar parcialmente el córtex visual o el córtex auditivo (dependiendo de la tarea que se esté realizando, y nunca buscando una desactivación total, pues es necesario mantener cierta alerta ante posibles mensajes de peligro), y esto implica un incremento de las oscilaciones  $\alpha$  en dicha zona. El incremento de la actividad  $\alpha$  en estos casos ayuda a prestar menos atención al mundo exterior, quizás aumentando la atención hacia el mundo interior, no sensorial. Cuando esto sucede, por supuesto, es más fácil cometer errores si se realiza una tarea repetitiva o aburrida. El análisis de la actividad  $\alpha$  registrada por un EEG en el córtex somatosensorial puede ser, por tanto, utilizado para predecir errores en la actuación de un sujeto y también, por supuesto, para confirmar si una persona determinada tiene o no problemas de atención, como el TDA.

Quizás las diferentes fases del sueño, que se caracterizan completamente en términos de la variación en las oscilaciones observadas en un EEG, conforman el ejemplo mejor conocido y más estudiado de cómo los distintos ritmos oscilatorios del cerebro están vinculados con un proceso neurofisiológico concreto.

#### FIGURA 3

EEG en las distintas fases del sueño. De arriba abajo: despierto, estados 1, 2 y 3 de la fase no-REM del sueño y, finalmente, fase REM o sueño paradójico.



Detallamos ahora algunos de los ritmos del cerebro cuyo significado es conocido y que pueden ser “entrenados” mediante *neurofeedback* o electroestimulación para mejorar diferentes tipos de anomalías:

- Ondas  $\delta$  (entre 1 y 4 Hz): Están asociadas al sueño. En bebés pueden ocupar el 40% de la actividad EEG, mientras que en adultos dicha proporción es del 5%, aproximadamente. Los niños con TDAH pueden mostrar un exceso de actividad  $\delta$ . Este ritmo puede entrenarse para su inhibición, pero nunca debe entrenarse su activación.
- Ondas  $\theta$  (entre 4 y 8 Hz): Este ritmo se asocia con la creatividad y la espontaneidad, pero también con la inatención, la facilidad para distraerse, el soñar despierto, la depresión y la ansiedad. Por ejemplo, un exceso de actividad  $\theta$  puede ser síntoma de depresión, ansiedad y otros trastornos emocionales. En personas con TDAH, suele entrenarse el descenso del cociente  $R_{\theta/\beta}$  o índice de Monastra (que definimos en el capítulo 2) precisamente inhibiendo la actividad  $\theta$  (y activando la

actividad  $\beta$ ).

- Ondas  $\alpha$  (entre 8 y 13 Hz): Se asocian con la meditación y la calma interior. Amplitudes elevadas en las zonas anteriores del encéfalo son usuales en niños que sueñan despiertos, y en personas con depresión. En casos de depresión y TDAH se observan amplitudes  $\alpha$  muy elevadas en las regiones anterior o frontal del encéfalo. Las amplitudes  $\alpha$  son normalmente más elevadas en las regiones posteriores que en las regiones anteriores. Sin embargo, lo normal es observar bastante simetría lateral (entre hemisferios izquierdo y derecho), aunque las amplitudes pueden ser inferiores en el hemisferio izquierdo (pero nunca más del 25%). Sin embargo, en personas con depresión se observan a menudo amplitudes  $\alpha$  más elevadas en el hemisferio izquierdo.
- Ondas SMR (entre 12 y 16 Hz): Este ritmo también se denomina  $\beta$ -bajo. Se manifiesta de forma predominante en el córtex somatosensorial y el córtex motor: electrodos C3, Cz y C4. Aumenta cuando los circuitos motores del cerebro están inactivos. De ahí que su amplitud aumenta en ausencia de movimiento y disminuye con el movimiento. Stermán y colaboradores fueron pioneros en entrenar este ritmo. Primero lo hicieron en gatos y, posteriormente, en humanos que padecían epilepsia. Concretamente, se entrena la activación de la banda 12-15 Hz en el electrodo Cz. Posteriormente, para el tratamiento de epilepsia y de hiperactividad, se ha combinado este protocolo con la inhibición de la actividad  $\theta$ . Este ritmo es también entrenado actualmente por deportistas de elite para mejorar su rendimiento.
- Ondas  $\beta$  (entre 12 y 21 Hz): Una actividad  $\beta$  excesiva la podemos encontrar en numerosos trastornos como, por ejemplo, en personas que padecen déficit de atención, trastorno obsesivo compulsivo, depresión o ansiedad. Es usual encontrar, en estos casos, que la actividad  $\beta$  anterior en el hemisferio derecho excede a la actividad  $\beta$  en el hemisferio izquierdo. En el caso de que se esté entrenando la activación de las

ondas  $\beta$ , hay que ser especialmente cuidadosos porque la banda  $\beta$  se superpone con las ondas detectadas en el EEG como consecuencia de actividad electromiográfica (es decir, los potenciales eléctricos generados por actividad muscular esquelética) y en ningún caso es aconsejable entrenar un artefacto del EEG.

- Ondas  $\beta$ -altas (20-32 Hz): Este ritmo está asociado al procesamiento cognitivo y los estados de preocupación y ansiedad, entre otros. A veces se encuentra presente simplemente porque el cerebro trata de compensar una actividad  $\theta$  excesiva. En *neurofeedback* lo usual es entrenar su inhibición.
- Ondas  $\gamma$  (entre 38 y 42 Hz): La actividad en torno a 40 Hz se suele detectar cuando el individuo se enfrenta a tareas complejas como la resolución de problemas. Frecuentemente, este ritmo no se produce (o apenas se produce) cuando hay problemas de aprendizaje o cuando existe una deficiencia mental. La actividad  $\gamma$  está presente en todos los electrodos del EEG, no existiendo ninguna región del encéfalo a la que se encuentre vinculada de forma especial. Se cree que este ritmo ayuda a que el cerebro se autoorganice y favorece el aprendizaje y la agudeza mental. Se activa cuando el cerebro necesita realizar alguna actividad especializada.

## Potenciales evocados

Sabemos que los potenciales eléctricos grabados por un EEG son extraordinariamente pequeños, del orden de microvoltios (un microvoltio son  $10^{-6}$  voltios). Además, en el proceso de grabación es muy sencillo que aparezca ruido, y dicho ruido pudiera ser de la misma amplitud que la señal que deseamos grabar y estudiar posteriormente. Entonces, ¿cómo podemos separar el ruido de la señal? ¿Cómo podemos, por ejemplo, diferenciar la respuesta eléctrica del encéfalo a un estímulo externo del ruido recogido por el sistema (o incluso, de la actividad eléctrica espontánea)? Además, si nuestro interés está relacionado con

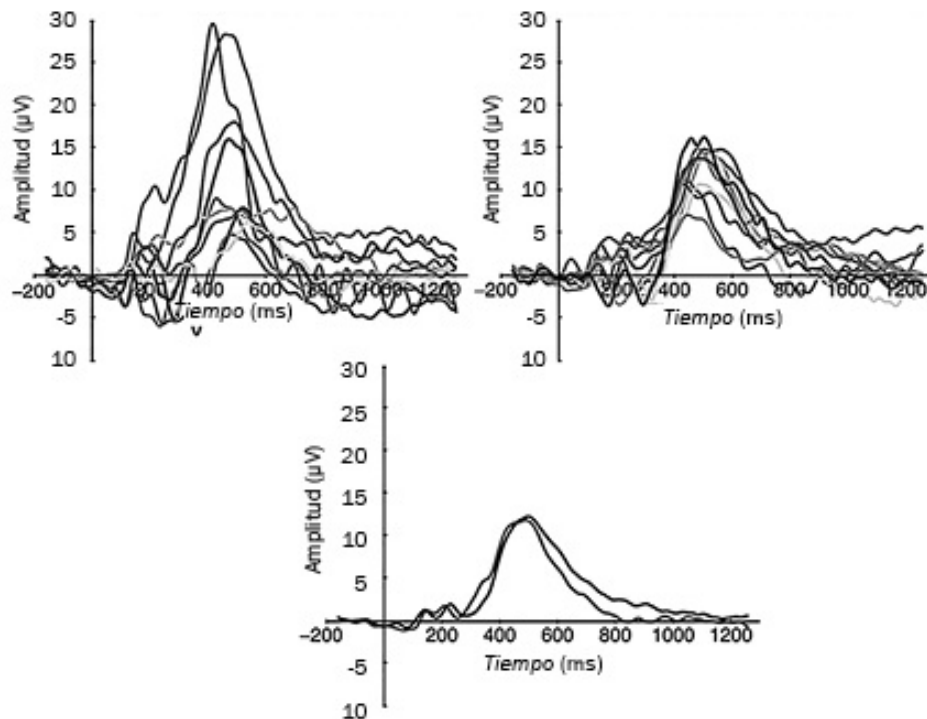
el estudio de las respuestas a estímulos (por ejemplo, queremos estudiar la respuesta atencional de un sujeto), entonces la actividad espontánea del cerebro (es decir, la actividad eléctrica que se produce en el encéfalo sin estar vinculada a ningún proceso neurocognitivo concreto), será recogida por el EEG, pero no nos interesa y, por tanto, debe interpretarse como ruido. ¿Cómo podemos eliminar este ruido de la grabación? En otras palabras, ¿cómo podemos registrar en el EEG solamente la actividad eléctrica que produce el encéfalo en respuesta a un determinado estímulo? Dicha actividad recibe el nombre de potencial evocado.

Una forma sencilla pero eficaz de lograrlo consiste en repetir el mismo estímulo numerosas veces y separar la información recogida por el sistema en intervalos de tiempo que comienzan un poco antes y finalizan un poco después de cada estímulo. Normalmente, se toman 100 milisegundos preestímulo y en torno a 500 milisegundos postestímulo, aunque estas cantidades podrían variar en función del tipo de estudio que se está realizando. Estos intervalos temporales reciben el nombre genérico de épocas. Posteriormente se analiza el EEG registrado y, cuando sea necesario, se eliminan los artefactos que contenga. Esto podría implicar la eliminación total de las épocas para las que se detecte que contienen artefactos, así como una modificación global del EEG para eliminar, por ejemplo, el ruido de la línea. A continuación, una vez disponemos de un EEG limpio, se calcula la media aritmética de las épocas asociadas a un estímulo-tipo concreto (hay diferentes tipos de estímulo, y solo se promedian los del mismo tipo). Con este sencillo mecanismo se logra eliminar el ruido aleatorio, por la sencilla razón de que este queda reducido a su esperanza, que es un valor constante. La curva resultante después de promediar es el potencial evocado (ERP, en siglas, del inglés *event related potential*). Evidentemente, las curvas obtenidas serán más suaves y más estables cuando el promedio se hace sobre un número elevado de épocas y cuando los datos se han obtenido con una frecuencia de muestreo suficientemente elevada. 100 épocas y 500 Hz bastarán para obtener resultados fiables. Además, cuando hemos fijado un experimento

determinado y lo aplicamos a un grupo numeroso de personas, conviene tomar el promedio de las curvas ERP sobre toda la población, pues las nuevas curvas ERP calculadas son mucho más estables que las que podemos hallar en cada uno de los sujetos. Este mecanismo se conoce como “gran promedio” y es muy usual en estudios en neurofisiología y neuropsicología. Para mostrar su utilidad, incluimos el ejemplo descrito por la figura 2. Veinte sujetos realizaron la misma tarea de tipo sensorial y se grabó el potencial evocado por dicha tarea en el mismo electrodo, en todos los sujetos, tomando épocas que comenzaban 200 ms antes y finalizaban 1.200 ms después del estímulo. El grupo de personas se dividió, al azar, en dos grupos de 10 individuos. Dibujamos el ERP de cada uno de ellos, agrupando las curvas en dos gráficas. En ambas se observa que existen importantes diferencias entre las curvas ERP individuales dentro de un mismo grupo de individuos. A continuación se toma el gran promedio asociado a cada uno de estos grupos, generando de esta forma dos curvas ERP que se dibujan conjuntamente en la tercera gráfica. Como se observa, la realización del gran promedio provoca que ambas curvas sean, ahora sí, bastante similares en amplitud, latencia y forma.

Se sabe, gracias a numerosos estudios fisiológicos basados en técnicas de neuroimagen, que tanto el procesado como la percepción sensorial de los estímulos externos se producen simultáneamente en varias regiones del córtex. Los potenciales evocados por estímulos sensoriales son generados por actividad sincronizada de las células piramidales que se encuentran orientadas perpendicularmente al cráneo, en dichas áreas del córtex. Resulta razonable, también por motivos fisiológicos relacionados con la denominada conducción de volumen, asumir que estos potenciales evocados reflejan la activación de ciertos potenciales eléctricos, espacialmente localizados, que se activan brevemente y de forma sincronizada en las distintas áreas estimuladas por la llegada de entradas sensoriales.

Estabilidad de los potenciales evocados obtenidos por gran promedio.



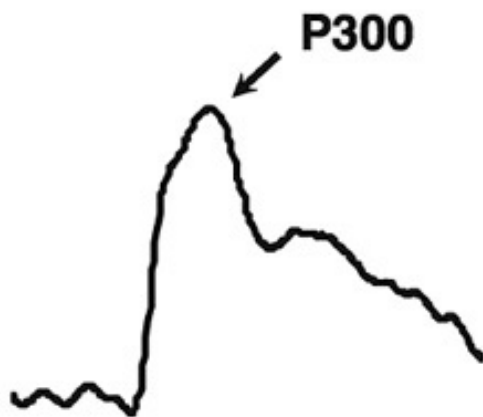
Existen numerosos estudios sobre potenciales evocados. La mayoría se centra en determinar las llamadas componentes del ERP. Cuando mostramos la curva descrita por un potencial evocado, en cualquiera de los electrodos del EEG, observamos una curva suave que pasa por varios picos y valles, que varían en amplitud, polaridad y duración. Un primer análisis de los potenciales evocados podría centrarse en la localización de estos picos, la determinación de sus polaridades, sus amplitudes, y sus latencias. En principio, cabría esperar que dichas características fuesen el resultado de diferentes procesos neuronales que, por supuesto, habría que identificar. Si bien es cierto que esta visión se ha corroborado en cierto grado por las numerosas investigaciones realizadas en el área, la verdad es que cada potencial evocado acaba siendo descrito como superposición de una o varias componentes —que tienen significado fisiológico— pero dichas componentes no se identifican de una forma inmediata y sencilla con la descripción de la onda en términos de picos y valles. Así, llamamos onda ERP a la curva calculada como promedio de las épocas asociadas a un mismo



estímulo tipo, siempre que se hayan tomado un número suficiente de estas (100 es aceptable). La onda ERP se asume que contiene información relacionada con los procesos sensoriales, cognitivos, afectivos, y motores provocados por un determinado estímulo. Un pico de la onda ERP es un extremo relativo de la onda que tiene, además, la característica añadida de ser confiable, en el sentido de que no se puede atribuir a la presencia de ruido de altas frecuencias en la señal. Por otra parte, cuando hablamos de componentes ERP nos estamos refiriendo a cambios en la onda ERP registrada que son atribuibles a procesos neuronales o psicológicos específicos, bien determinados. Evidentemente, esta definición no es sencilla de aplicar porque su uso conlleva la determinación causal que liga un proceso neurocognitivo concreto —como, por ejemplo, una toma de decisión— con la aparición de cierto rasgo característico en la onda ERP, y esto requiere un arduo trabajo de investigación previo, que implica la colaboración de un amplio equipo multidisciplinar. Aun así, existen infinidad de estudios que han logrado identificar ciertas componentes ERP, definiendo además los mecanismos apropiados para su cálculo. Cuando una componente ERP es aceptada como tal, se bautiza y se utiliza como posible neuromarcador de diferentes patologías. Quizás la componente ERP más estudiada es la P300, aunque hay muchas otras, como las componentes MMN, N1, N170 o LRP. La componente P300 fue descubierta por Sutton y colaboradores en 1965 y sirvió para estimular fuertemente el uso de técnicas de neurometría para el estudio de procesos neurocognitivos. Aparece en los electrodos Fz, Cz y Pz cuando sometemos a un sujeto a un estímulo sensorial (que puede ser visual o auditivo) poco frecuente, bien de forma aislada o bien inmerso entre otros estímulos frecuentes, y se le pide que reaccione a dicho estímulo pulsando un botón. Cuando realizamos esta experiencia repetidas veces y calculamos, en los correspondientes nodos, la curva ERP, aparece una curva característica, que denominamos componente P300, cuyo primer máximo confiable posee una amplitud muy superior a los primeros extremos relativos de la onda y que se produce en torno a 300 milisegundos posteriores a la emisión del estímulo infrecuente. Los datos más

relevantes de esta componente son su latencia (es decir, el tiempo exacto transcurrido entre la emisión del estímulo y la aparición del primer máximo confiable), su amplitud (es decir, la amplitud del correspondiente máximo) y su forma. La latencia y la amplitud se guardan, para su posterior estudio. Se asume que dichos valores se comportan como los valores que se obtienen al realizar un experimento gobernado por una variable aleatoria (v.a. en lo sucesivo) normal y, por tanto, cuando se ha fijado la media y la desviación típica de los mismos para un grupo de pacientes sanos, estos se pueden utilizar como neuromarcadores de una posible patología. La componente P300 está asociada a la atención y la memoria.

FIGURA 5  
Componente P300.



No es sencillo determinar una nueva componente ERP, pero cuando esto se logra siempre supone un avance significativo en las investigaciones en neuropsicología y psiquiatría. Las primeras componentes ERP se detectaron, sin la necesidad de hacer promedios, en la década de 1930, pero la llegada de los computadores, entre otros avances tecnológicos, permitieron introducir en la década de 1960 la técnica de los promedios, y esto supuso en verdadero impulso en el área, especialmente a partir de principios de 1980.

## CAPÍTULO 2

# Todo son ondas: una breve incursión en el análisis de Fourier

Pensemos, por un momento, en la señal recogida por cada uno de los electrodos de un electroencefalograma. Esta señal es, evidentemente, unidimensional y varía con el tiempo. El electrodo ha tomado muestras de la actividad eléctrica en ese preciso punto del cuero cabelludo donde está ubicado. Cada segundo hemos registrado entre 250 y 2.000 muestras —depende del aparato que usemos para grabar, por supuesto—. Esta es la actividad eléctrica básica que luego debemos analizar. ¿Por qué tomamos muestras y por qué se toman precisamente con esta frecuencia de muestreo? Y si tomamos muestras, ¿cómo sabemos que no vamos a perder información relevante? La respuesta a estas preguntas está en el análisis de Fourier y fue generada poco a poco por varios científicos, con intereses muy diversos, entre mediados del siglo XIX y mediados del siglo XX, involucrando algunos de los nombres fundamentales para la comprensión del desarrollo de la matemática moderna y, además, está conectada con algunas cuestiones profundas de la física y la ingeniería.

Pero antes de entrar en materia consideramos necesario introducir algunos conceptos y fórmulas que serán posteriormente fundamentales para comprender las ideas y resultados que pretendemos relatar. Para empezar, recordemos que un número complejo es sencillamente una expresión del tipo  $a+ib$ , donde  $i$  denota la “unidad imaginaria” (esto es, resuelve la ecuación  $i^2+1=0$ ). El par  $(a,b)$  representa unívocamente al número  $a+ib$  (son sus “coordenadas cartesianas”). Para números complejos no nulos hay otra representación natural que, para nosotros, será de enorme importancia. En efecto, si dibujamos el vector  $z=(a,b)=$

$a+ib$  en el plano, veremos que este queda determinado a partir de su módulo  $\rho = |z| = \sqrt{a^2 + b^2}$  y del ángulo  $\theta$  que forma con el eje de abscisas. De hecho,  $z = |z|(\cos\theta + i\sin\theta)$ . Esta es la forma polar del número complejo  $z$ .

La función  $E(\theta) = \cos\theta + i\sin\theta$  es especialmente importante. Satisface la ecuación diferencial  $E'(\theta) = iE(\theta)$ , con condición inicial  $E(0) = 1$ , ecuación que caracteriza la función exponencial  $e^{i\theta}$ . Por tanto,  $e^{i\theta} = \cos\theta + i\sin\theta$  y la forma polar del número complejo  $z$  queda  $z = \rho e^{i\theta}$ . Esta expresión aparece constantemente en nuestros cálculos y, por tanto, conviene familiarizarse con ella. En particular, cuando estudiamos el comportamiento en frecuencias de una señal, nos interesa más la forma polar de cada componente que la expresión en coordenadas cartesianas, porque tiene un significado físico más nítido. La energía de cada componente queda determinada a partir de su módulo (es el cuadrado de dicho valor). Por otra parte, la fase nos habla de la temporalización (estamos interesados en procesos que son cíclicos —oscilaciones— y la fase es un ángulo, por lo que podemos identificarla con la medida de un reloj). Se sabe que la fase sirve para que las neuronas determinen el momento en el que van a activar sus potenciales de acción, mientras que la energía detectada refleja si se ha activado un mayor (o menor) número de neuronas. Ambas cuestiones son vitales para que el cerebro responda adecuadamente a los estímulos que recibe y son fundamentales, por ejemplo, en los estudios relacionados con las habilidades de atención de las personas.

## Vivir entre dos mundos: tiempo y frecuencia

Para empezar, hay que saber que “cualquier” señal periódica se puede representar como superposición de ciertas señales básicas, llamadas armónicos elementales. Hemos escrito entre comillas la palabra “cualquier” porque en realidad la afirmación no se satisface para toda señal periódica pero, como veremos, sí para todas las que nos puedan interesar en las aplicaciones físicas. Esto no fue fácil de intuir y, mucho menos, demostrar. La intuición viene de

tiempos lejanos, tan lejos como Pitágoras, quien se ocupó ya de cuestiones musicales e introdujo la idea del armónico puro. La acústica sería, de hecho, la principal motivación para D. Bernoulli en sus intentos de justificar su estudio de la ecuación de ondas, asociada al problema de la cuerda vibrante. Por su parte, es seguro que Fourier conocía los trabajos de Bernoulli y, con toda probabilidad, su fuerte convicción de que toda función se puede expresar como superposición de ondas sinusoidales simples tuvo también bases físicas. No olvidemos que, después de todo, Fourier es considerado uno de los fundadores de la física moderna por sus contribuciones al estudio de la propagación del calor, siendo su aportación más importante la creación de un modelo matemático, basado en ecuaciones en derivadas parciales, que describe cómo varía la distribución del calor, en función del tiempo y la posición, en un sólido conductor, y la introducción del método de separación de variables para la solución de las ecuaciones implicadas. Dicho método implicaba la hipótesis, que defendió con enorme entusiasmo durante toda su vida, de que las funciones “arbitrarias” se pueden descomponer como sumas (o superposiciones, si empleamos una terminología más física) de ciertas funciones simples: los armónicos puros, u ondas básicas. Aunque Fourier erró en algunos de los (numerosos) argumentos que utilizó para defender sus tesis, y fue duramente criticado por ello, su intuición fue finalmente validada en su más profunda esencia y, como suele suceder en matemáticas, el establecimiento definitivo de los límites que esta alcanza, fue tema de investigación, no por unos años o unas décadas, sino durante siglos, llegando a delimitarse en sus últimos aspectos solo a mediados de la década de 1960, con el trabajo de L. Carleson, a quien se le concedería el Premio Abel en 2006, en buena medida por estos logros. Además, el modelo de “descomponer una función arbitraria como suma de ciertas funciones básicas” fue seguido en diferentes contextos del análisis, dando lugar a una forma de pensar en matemáticas que, en definitiva, puede verse como uno de los gérmenes de, por ejemplo, el análisis funcional.

## Series de Fourier

Bajando a la tierra, y para explicar con mayor concreción a qué tipo de resultados nos estamos refiriendo, vamos a enunciar algunos de ellos. Si tenemos una función mínimamente regular y  $T$ -periódica  $f(t)$  (es decir,  $f(t)=f(t+T)$  para todo número real  $t$ ) entonces esta función se puede describir como una suma (posiblemente infinita) de armónicos puros. Es decir,

$$f(t) = \frac{a_0}{2} + \sum_{k=1}^{\infty} \left( a_k \cos \frac{2\pi kt}{T} + b_k \sin \frac{2\pi kt}{T} \right)$$

para ciertos coeficientes  $a_k, b_k$ , que dependen exclusivamente de  $f$  y que reciben el nombre de coeficientes de Fourier de la función.

Prestemos atención a la ecuación anterior. En el miembro de la izquierda tenemos una función  $f(t)$  que depende del tiempo —pensemos que  $t$  representa el paso del tiempo y que  $f(t)$  es la magnitud de la señal eléctrica que medimos en el instante de tiempo  $t$ — mientras que en el miembro de la derecha de la ecuación, aparecen ciertas funciones fijas (los senos y cosenos) que se superponen (es decir, se suman) y se promedian (es decir, se multiplican por ciertos números).

¿Qué puede significar que estas funciones estén fijas? Ellas no dependen de la forma de la señal  $f(t)$ : usamos las mismas funciones  $\cos \frac{2\pi kt}{T}$ ,  $\sin \frac{2\pi kt}{T}$ , para toda función periódica  $f(t)$  con periodo  $T$ . Si pintamos una de estas funciones, veremos que se trata de una onda que oscila de forma regular, repitiendo el mismo trazado  $k$  veces en cada intervalo de longitud  $T$ . Es decir,  $k$  nos da la frecuencia con la que esta función oscila en cada intervalo periodo. Los coeficientes  $a_k, b_k$  nos dicen cuánta energía aporta esta vibración fundamental a nuestra señal  $f(t)$ . Así, son los coeficientes los que determinan completamente a  $f(t)$  y, además, su significado se debe interpretar en términos de frecuencias. A un lado de la ecuación tenemos una descripción de  $f(t)$  en función del tiempo y al otro lado, recuperamos la misma señal, vista desde un universo diferente: el mundo de las frecuencias. Lo más importante en esto es que podemos viajar de

un lado al otro sin perder información. Este fue el maravilloso descubrimiento de Fourier, que daría lugar a importantes cambios no solo en análisis matemático sino también en física, ingeniería, y numerosos campos de aplicación, incluyendo la medicina.

Dada la función  $T$ -periódica  $f(t)$ , sus coeficientes de Fourier se pueden calcular explícitamente gracias a las fórmulas de Fourier:

$a_k = \frac{1}{T} \int_0^T f(t) \cos \frac{2\pi kt}{T} dt$  y  $b_k = \frac{1}{T} \int_0^T f(t) \sin \frac{2\pi kt}{T} dt$ . En particular, estos pueden obtenerse tan pronto como la función  $f$  sea integrable en  $[0, T]$ . En efecto, asumamos momentáneamente que la función  $f$  admite una descomposición del tipo anterior, y que la convergencia de la serie que aparece en el segundo miembro de la ecuación es tan fuerte que permite multiplicar a ambos lados por cualquier función suficientemente regular (como, por ejemplo, por las funciones  $\cos \frac{2\pi kt}{T}$  o  $\sin \frac{2\pi kt}{T}$ ) e integrar término a término, conservando la igualdad. Si hacemos esto y tenemos, además, en consideración las identidades

$$\frac{1}{T} \int_0^T \cos \frac{2\pi nt}{T} \cos \frac{2\pi mt}{T} dt = \delta_{n,m},$$

y

$$\frac{1}{T} \int_0^T \sin \frac{2\pi nt}{T} \sin \frac{2\pi mt}{T} dt = \delta_{n,m},$$

para todo  $n, m \in \mathbb{Z}$  (donde  $\delta_{n,m}$  se anula si  $n \neq m$  y vale 1 si  $n=m$ ), entonces es evidente que los coeficientes de Fourier  $a_k, b_k$  deben satisfacer las fórmulas de Fourier. Estos cálculos son en realidad una motivación (poderosa, pero meramente heurística) para aceptar las fórmulas de Fourier como las ecuaciones apropiadas que deben definir los coeficientes de Fourier de la función  $f$ , pero de ningún modo hemos demostrado que estas ecuaciones se dan para toda función  $f$ , pues ello nos llevaría al círculo vicioso de asumir, de partida, unas propiedades de convergencia de las series de Fourier que, en realidad, están en el aire.

De todas formas, la solución a este problema es sencilla: forzamos las fórmulas de Fourier como las apropiadas para definir la serie de Fourier de la función  $f$ , pero esta asociación se asume como algo meramente formal. Una vez esto ha quedado claro, se plantea la cuestión de saber bajo qué condiciones es convergente la serie de Fourier que hemos asociado formalmente a  $f$  y, caso de serlo, si converge a la función de partida, y en qué sentido.

Resulta sorprendente, y un descubrimiento maravilloso, saber que las condiciones que garantizan la convergencia de las series de Fourier a la función original son extraordinariamente débiles.

El primer resultado que se demostró rigurosamente en esta dirección fue publicado en 1829, se debe a Dirichlet y zanjó de forma contundente, a favor de Fourier, la polémica que se había generado en torno a sus ideas, durante años y años de desagradables conflictos (en los que no solo los aspectos científicos fueron considerados, sino también cuestiones personales y políticas). Dirichlet demostró que si  $f$  es derivable a trozos y con derivada continua a trozos, entonces su serie Fourier converge a la propia función  $f$  uniformemente sobre compactos contenidos en la zona de continuidad de  $f$  y, en los puntos de discontinuidad, converge al punto medio del salto. Por supuesto, estos resultados se fueron refinando con el tiempo, y dieron lugar a una importante rama del análisis matemático que no ha parado de crecer y producir nuevas ideas y que ha sido fundamental en numerosas aplicaciones de la matemática.

Aunque a primera vista podría ser difícil de adivinar, el cálculo que hemos improvisado para los coeficientes de Fourier se basa, en realidad, en un argumento geométrico. Lo que esconde es, por así decirlo, la introducción de una idea de “espacio” que facilita el arte de medir distancias en un contexto donde existen infinitas dimensiones. Cada coeficiente de Fourier, que explica el comportamiento de la señal en una frecuencia determinada, representa una dimensión del espacio en cuestión.

El concepto de espacio que todos hemos incorporado de un modo más o menos intuitivo a nuestra experiencia sensorial del mundo, no es otro que el



espacio Euclídeo de dimensión tres,  $\mathbb{R}^3$ . En él podemos medir distancias y calcular ángulos de forma sencilla. Gracias a la aparición del álgebra, estas operaciones se pudieron sistematizar en torno a un concepto sencillo, el producto escalar de vectores. Dados los vectores  $x=(x_0,x_1,x_2)$  e  $y=(y_0,y_1,y_2)$  del espacio

ordinario, su producto escalar viene dado por  $x,y = \sum_{i=0}^2 x_i y_i$ . La distancia que los

separa está entonces dada por  $d(x,y) = \|\hat{x}-\hat{y}\| = \sqrt{\langle x-y, x-y \rangle}$  y el ángulo que forman proviene de despejar  $\theta$  en la fórmula  $\langle x,y \rangle = \|\hat{x}\| \|\hat{y}\| \cos\theta$ . En particular, los vectores serán perpendiculares si su producto escalar vale 0.

Supongamos por un momento que tomamos dos vectores  $u,v$  del plano  $\mathbb{R}^2$ , que son perpendiculares. Entonces cualquier otro vector  $x \in \mathbb{R}^2$  se puede escribir de forma única como  $x = au + bv$  para ciertos valores  $a,b \in \mathbb{R}$ . Es más, dichos valores se obtienen fácilmente a partir del producto escalar gracias al cálculo siguiente:

$$\begin{aligned} \langle x,u \rangle &= \langle au + bv, u \rangle \\ &= a\langle u,u \rangle + b\langle v,u \rangle \\ &= a\langle u,u \rangle, \end{aligned}$$

ya que  $\langle v,u \rangle = 0$  al ser dichos vectores perpendiculares. Esto nos conduce a

concluir que  $a = \frac{\langle x,u \rangle}{\langle u,u \rangle}$ . Análogamente,  $b = \frac{\langle x,v \rangle}{\langle v,v \rangle}$ .

Estas ideas se pueden extender de forma inmediata al espacio Euclídeo de dimensión  $N$ ,

$$\mathbb{R}^N = \{x = (x_0, \dots, x_{N-1}) : x_i \in \mathbb{R}, i = 0, \dots, N-1\}$$

y, lo que es más aún, se pueden trasladar a un contexto mucho más amplio simplemente haciendo abstracción de los conceptos de “espacio” y “producto

escalar en un espacio”.

Por “espacio” entendemos cualquier conjunto en el que se han introducido las operaciones de suma y producto por un escalar, y en el que estas satisfacen las propiedades usuales de los vectores del espacio ordinario. Estos espacios representan una buena aproximación al comportamiento de numerosos procesos físicos en los que se producen fenómenos de superposición y modulación (atenuación o amplificación) de magnitudes. Por ejemplo, el sonido viaja en ondas que, cuando se encuentran, se suman. Cuando escuchamos varios sonidos simultáneamente lo que estamos oyendo es, en realidad, la superposición o suma de los mismos. Además, un sonido será más “fuerte” o más “débil” (o, si se quiere, se percibirá como más cercano o más lejano) en función de la intensidad de la onda —no de su forma, ni de su frecuencia de vibración—, y dicha intensidad cambia cuando multiplicamos por un número toda la señal sonora. Por tanto, las ondas sonoras representan perfectamente los elementos de un “espacio” en el sentido que acabamos de describir. Lo mismo sucede con las ondas electromagnéticas y con las señales eléctricas recogidas en un EEG. El nombre que se usa en matemáticas para este concepto de “espacio” es “espacio vectorial” y se le pone el apellido “real” cuando los coeficientes por los que multiplicamos los vectores son números reales y “complejo” cuando dichos coeficientes son números complejos.

Si  $H$  es un espacio vectorial real, decimos que la aplicación  $\cdot : H \times H \rightarrow \mathbb{R}$ , que lleva un par de vectores  $x, y \in H$  al número real  $\langle x, y \rangle$ , es un producto escalar en  $H$ , si es una aplicación simétrica (es decir,  $\langle x, y \rangle = \langle y, x \rangle$ ), es lineal en cada una de sus variables (para lo cual basta asumir que,  $\langle ax + by, z \rangle = a \langle x, z \rangle + b \langle y, z \rangle$  para todo  $x, y, z$  y todo  $a, b \in \mathbb{R}$ , pues la linealidad en la otra variable se obtiene gratis a partir de la propiedad de simetría anterior) y, por último, es una aplicación definida positiva (es decir,  $\langle x, x \rangle \geq 0$  cuando  $x \in H$  y  $\langle x, x \rangle = 0$  solo si  $x = 0$ ). Los espacios vectoriales sobre los que se ha definido un producto escalar se llaman Euclideos porque en ellos es posible introducir los conceptos básicos

de la geometría Euclídea: distancias y ángulos. Si  $H$  es un espacio vectorial complejo, existe un concepto similar de producto escalar. Lo único que tenemos que modificar es que los números  $\langle x, y \rangle$  con  $x, y \in H$  deben ser ahora números complejos y que la simetría es sustituida por la propiedad hermitica:

$\overline{\langle x, y \rangle} = \langle y, x \rangle$ . Por ejemplo, en  $\mathbb{C}^N$  el producto escalar usual está dado por

$$\langle x, y \rangle = \sum_{i=0}^{N-1} x_i \overline{y_i} \quad L^2(0, T) = \left\{ f : \int_0^T |f(t)|^2 dt < +\infty \right\}$$

. En el caso de , el espacio de las señales complejas  $T$ -periódicas de energía finita, se considera el producto escalar

$$\langle f, g \rangle = \int_0^T f(t) \overline{g(t)} dt \quad d(f, g) = \left( \int_0^T |f(t) - g(t)|^2 dt \right)^{\frac{1}{2}}$$

Los coeficientes de Fourier de la función  $T$ -periódica  $f(t)$  se calculan directamente a partir de este producto escalar, pues

$$a_k = \frac{\langle f, \cos \frac{2\pi kt}{T} \rangle}{\langle \cos \frac{2\pi kt}{T}, \cos \frac{2\pi kt}{T} \rangle} \quad b_k = \frac{\langle f, \sin \frac{2\pi kt}{T} \rangle}{\langle \sin \frac{2\pi kt}{T}, \sin \frac{2\pi kt}{T} \rangle} \quad \text{y}$$

Además, si tenemos en cuenta que  $e^{i\theta} = \cos \theta + i \sin \theta$  y, por tanto,

$$\cos \theta = \frac{e^{i\theta} + e^{-i\theta}}{2} \quad \text{y} \quad \sin \theta = \frac{e^{i\theta} - e^{-i\theta}}{2i}$$

, resulta que la serie de Fourier de  $f$  admite una representación del tipo  $\sum_{k=-\infty}^{\infty} c_k e^{\frac{2\pi i kt}{T}}$ , donde los coeficientes de Fourier complejos  $c_k$  están dados por

$$c_k = \frac{\langle f, e^{\frac{2\pi i kt}{T}} \rangle}{\langle e^{\frac{2\pi i kt}{T}}, e^{\frac{2\pi i kt}{T}} \rangle} = a_k - ib_k$$

El espacio  $L^2(0, T)$ , cuando la integral que se considera es la de Lebesgue, tiene la bonita propiedad de que permite una identificación completa entre la

señal  $f \in L^2(0, T)$  y sus coeficientes de Fourier.

En general, si  $H$  es un espacio dotado de un producto escalar y  $\{\phi_k\}_{k=0}^\infty$  es una familia de elementos (no nulos) de  $H$  que genera un subespacio denso y que satisface las condiciones de ortogonalidad  $\langle \phi_n, \phi_m \rangle = 0$  cuando  $n \neq m$ , entonces se puede demostrar que todo elemento  $x \in H$  admite una expresión del tipo

$$x = \sum_{k=0}^{\infty} \frac{\langle x, \phi_k \rangle}{\langle \phi_k, \phi_k \rangle} \phi_k$$
, siendo la convergencia en el sentido de la distancia definida en  $H$ . Este resultado generaliza, evidentemente, las series de Fourier a un contexto mucho más amplio, y ha sido utilizado de forma muy intensa en diferentes ramas de la matemática y de la física. Los coeficientes de Fourier del vector  $x$  respecto

de la base ortogonal son, entonces, los números 
$$c_k(x) = \frac{\langle x, \phi_k \rangle}{\langle \phi_k, \phi_k \rangle}, \quad k=0,1,2,\dots,$$
 y conocer dichos coeficientes nos permite una recuperación completa del vector gracias a la fórmula anterior.

Un caso extremo, en el que la cuestión de la convergencia de las series de Fourier generalizadas desaparece, es cuando nuestro espacio tiene dimensión finita. Concretamente, podemos asumir que trabajamos en  $\mathbb{C}^N$  con el producto

escalar usual, 
$$\langle x, y \rangle = \sum_{k=0}^{N-1} x_k \overline{y_k}.$$
 En este contexto hay dos bases ortogonales que son especialmente importantes. La primera es la base canónica, formada por los vectores  $e_k = (0, 0, \dots, 0, 1, 0, \dots, 0)$  (donde el 1 aparece en la posición  $(k+1)$ -ésima del vector). Cuando escribimos  $x = (x_0, x_1, \dots, x_{N-1})$  lo que estamos diciendo es que

$$x = \sum_{k=0}^{N-1} x_k e_k.$$
 Es decir, esta base nos proporciona de forma natural la representación de la señal  $x$  en el dominio del tiempo. El número  $x_k$  es sencillamente el valor de la señal en el instante  $k$  y se suele denotar, indistintamente, por  $x_k$  o  $x(k)$ . La

segunda base fundamental, es, evidentemente, la que nos describe estas señales en términos de frecuencias. Se trata de la base formada por los vectores

$E_k = \left( e^{\frac{2\pi i k n}{N}} \right)_{n=0}^{N-1}$ . Es decir, la componente  $n$ -ésima de  $E_k$  es el número complejo  $e^{\frac{2\pi i k n}{N}}$ , para  $n=0,1,\dots,N-1$ . Que  $\langle E_k, E_r \rangle = 0$  para  $k \neq r$ , por lo que estos vectores forman un sistema ortogonal en  $\mathbb{C}^N$ , es una sencilla cuenta, basada en aplicar a

$x = e^{\frac{2\pi i (k-r)n}{N}}$  la fórmula  $1 + x + \dots + x^{N-1} = \frac{x^N - 1}{x - 1}$ , la cual se demuestra simplemente multiplicando a ambos lados de la igualdad por  $x-1$ .

Se sigue que todo vector  $x \in \mathbb{C}^N$  admite una expresión del tipo

$$x = \sum_{k=0}^{N-1} \frac{x, E_k}{E_k, E_k} E_k = \frac{1}{N} \sum_{k=0}^{N-1} \hat{x}(k) E_k$$

donde

$$\hat{x}(k) = x, E_k = \sum_{n=0}^{N-1} x(n) e^{\frac{-2\pi i k n}{N}}$$

.

Llamamos transformada discreta de Fourier (o dft) del vector  $x$  al vector  $dft(x) = \hat{x} = (\hat{x}(0), \dots, \hat{x}(N-1))$ .

Del mismo modo que los coeficientes de Fourier de la señal periódica  $f(t)$  nos describen el comportamiento en frecuencias de esta, los coeficientes de la dft del

vector  $x$  nos dan (salvo un factor de  $\frac{1}{N}$ ) las coordenadas del vector  $x$  en la base de las frecuencias  $\{E_k\}_{k=0}^{N-1}$  de  $\mathbb{C}^N$ .

## La transformada de Fourier

Sabemos que la serie de Fourier de una función periódica converge con gran

facilidad a la función. Esto permite utilizar los desarrollos en serie de Fourier para aplicaciones en muy diversos contextos. Por supuesto, existen funciones cuya serie de Fourier no converge en uno u otro sentido, y la descripción precisa de los límites en torno a los cuales existen (o no) dificultades, es una teoría matemática muy sutil y elaborada, con enorme dificultad técnica, que escapa a lo que razonablemente podemos contar en este texto. Lo importante aquí es saber que las series de Fourier convergen en todos los ejemplos de funciones periódicas que podamos encontrar en el mundo real. Hay sin embargo una propiedad indispensable en todo este desarrollo, a la cual no hemos dado aún la importancia que tiene, por las limitaciones que impone en las aplicaciones, y es que las funciones bajo consideración han sido, hasta ahora, necesariamente periódicas. ¡Y los electroencefalogramas no son periódicos!

El propio Fourier fue consciente de las limitaciones que implica la hipótesis de periodicidad. De hecho, en su *Teoría Analítica del Calor*, cuando se ocupó del problema de la distribución del calor en sólidos infinitos, necesitó generalizar sus resultados al caso no periódico, para lo cual introdujo el operador que ahora llamamos “transformada de Fourier”. Si bien los coeficientes de Fourier no se pueden atribuir a él, la transformada o integral de Fourier sí se la debemos, sin ningún género de dudas.

Existen varias formas de introducir este operador. La más natural, y la que usó el propio Fourier, es la siguiente: Asumimos que nuestra función  $f(t)$  no es periódica, pero decae a 0 con cierta velocidad cuando  $t$  tiende a infinito. Podemos, por ejemplo, suponer que se trata de una señal definida en toda la recta

real y cuyo valor absoluto es sumable, es decir, tal que  $\int_{-\infty}^{\infty} |f(t)| dt < \infty$ .

Elegimos un valor  $T$  y consideramos la extensión  $T$ -periódica de  $f|_{[-T/2, T/2]}$ , que denotamos por  $x_T(t)$ . A continuación, nos fijamos en el  $k$ -ésimo coeficiente de

Fourier complejo de  $x_T(t)$ , que está dado por  $c_k(x_T) = \frac{1}{T} \int_{-T/2}^{T/2} f(t) e^{\frac{-2\pi i k t}{T}} dt$ . Si

hacemos ahora que  $T$  tienda a infinito, entonces

$$Tc_k(x_T) = \int_{-T/2}^{T/2} f(t) e^{-2\pi i k t} dt \cong \int_{-\infty}^{\infty} f(t) e^{-2\pi i k t} dt = \mathcal{F}(f)\left(\frac{k}{T}\right),$$

donde

$$\mathcal{F}(f)(u) = \int_{-\infty}^{\infty} f(t) e^{-2\pi i t u} dt$$

es la transformada de Fourier de la señal  $f(t)$ . Es decir, para  $T$  grande,  $Tc_k(x_T)$  es

una buena aproximación de  $\mathcal{F}(f)\left(\frac{k}{T}\right)$ . Ahora bien, los valores  $c_k(x_T)$  representan en términos frecuenciales a  $f(t)$  en  $[-T/2, T/2]$ . Por tanto, el efecto que tiene perder la periodicidad de la señal es que podemos variar libremente los valores de  $k$  y  $T$  para considerar frecuencias arbitrarias: el espectro pasa de ser discreto a ser todo el continuo de la recta real, pero mantiene su significado.

Por supuesto, una vez hemos definido el operador  $\mathcal{F}$ , es necesario demostrar que tiene las propiedades más importantes que tenían los coeficientes de Fourier para el caso de las señales periódicas. En particular, hay que probar que el viaje desde el mundo del tiempo al mundo de las frecuencias es de ida y vuelta y que, además, en los trayectos no se produce pérdida de información. Este es de hecho el contenido del llamado teorema integral de Fourier, que nos dice que, bajo hipótesis mínimas, el operador

$$\mathcal{F}^{-1}(f)(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} f(u) e^{2\pi i t u} du$$

actúa como inverso del operador  $\mathcal{F}$ . Hay que pensar en este resultado como un teorema sobre sumabilidad, parecido a los que garantizan que la serie de Fourier de una función periódica converge a dicha función.

La transformada de Fourier está definida, en principio, para funciones sumables, es decir, para señales con integral finita  $\int_{-\infty}^{\infty} |f(t)| dt < \infty$ . Además, esta definición se puede extender a todas las señales de energía finita.

También es importante saber que la transformada de Fourier se puede definir para el espacio de las “distribuciones temperadas”. Las distribuciones o funciones generalizadas son un ente abstracto que permite ampliar el concepto de función para facilitar el estudio de ciertos fenómenos físicos que no se pueden modelar con funciones ordinarias. En este texto no podemos introducir los formalismos necesarios para dicha teoría, pues son excesivamente complejos y sutiles, pero sí queremos enfatizar su importancia y su necesidad, incluso para aplicaciones que, en último término, se refieren exclusivamente a funciones ordinarias, algunas de cuyas propiedades se obtienen solamente a partir de su interpretación distribucional. Existen, además, numerosos problemas clásicos, como el problema de la cuerda vibrante, o el estudio de ciertas partes de la mecánica cuántica, cuya clarificación pasa necesariamente por el uso de distribuciones.

## Convolución

Un principio matemático fundamental para las aplicaciones del análisis de Fourier es el llamado teorema de convolución. El proceso de convolucionar dos señales es la forma matemática con la que se describen los llamados filtros de ondas. Desde el punto de vista físico es fundamental disponer de dispositivos (como, por ejemplo, circuitos eléctricos) que filtren las señales eliminando ciertas frecuencias y dejando pasar, sin tocarlas, otras. Por ejemplo, en el caso de los electroencefalogramas se sabe que la propia corriente eléctrica (los cables) introduce un ruido en la señal que se puede identificar usualmente como un armónico puro a 50 o 60 Hz, dependiendo de si estamos en Europa o en Estados Unidos. Además, típicamente, las señales EEG que nos interesan vibran a frecuencias inferiores a los 50 Hz. Sería, por tanto, deseable, disponer de un mecanismo que elimine las frecuencias de 50Hz y 60Hz (es decir, el ruido) del EEG.

Si miramos la transformada de Fourier de la señal  $f(t)$ , la mejor forma de permitir el paso de ciertas frecuencias y eliminar otras es multiplicar la



transformada  $\mathcal{F}(f)(u)$  por una función  $H(u)$  que tome valor 1 en aquellas frecuencias  $u$  cuyo paso queremos permitir y se anule en las frecuencias que deseamos eliminar. Si la función  $H(u)$  es la transformada de Fourier de cierta señal  $h$ , el teorema de convolución nos dice que la acción del filtro se describe, en el dominio del tiempo, mediante la operación de convolucionar  $f$  contra la función  $h$ . En fórmulas,

$$\mathcal{F}(f * h)(u) = \mathcal{F}(f)(u) \cdot \mathcal{F}(h)(u) = \mathcal{F}(f)(u) \cdot H(u),$$

donde

$$(f * h)(t) = \int_{-\infty}^{\infty} f(s)h(t-s)ds$$

es la operación de convolución.

## El teorema del muestreo

Este resultado nos proporciona la seguridad de una transición suave, sin pérdida de información, desde el mundo analógico al mundo digital. La razón por la que este principio es siempre aplicable es tan sencilla como que, del mismo modo que existen limitaciones para el tamaño de las partículas (la materia está cuantizada), también su capacidad de vibrar está sometida a límites específicos. Toda la información que recibimos del mundo exterior a través de nuestros sentidos, por vía de magnitudes físicas como la luz, el sonido, la presión o el calor se puede interpretar como resultado de la vibración de ciertas partículas. Además, nuestro sistema nervioso, nuestros sentidos, imponen una restricción aún mayor, si cabe, a las vibraciones que podemos percibir. Por tanto, en todos los casos de interés, las señales registradas por nuestros sentidos, o por cualquier sistema de medición del que dispongamos, son necesariamente de banda limitada. Por ejemplo, sabemos que el oído humano es incapaz de percibir sonidos que se emiten con una frecuencia inferior a los 20 Hz o superior a 20.000 Hz. Por tanto, todas las señales sonoras que podemos percibir están amparadas por el teorema del muestreo. Esto permite crear aparatos de

reproducción digitales como los reproductores de CD, en los que no existe pérdida de calidad del sonido respecto de los sonidos percibidos. Por supuesto, las señales eléctricas que grabamos en un EEG también están limitadas en banda, por lo que podemos muestrearlas sin perder información, y podemos tratarlas digitalmente.

¿En qué se basa el teorema del muestreo? Supongamos, para empezar, que tenemos una señal periódica. Diremos que la señal está limitada en banda si para su descripción solo se requiere utilizar un conjunto finito de frecuencias o, lo que es lo mismo, si todos sus coeficientes de Fourier, excepto unos pocos (un número finito de ellos), se anulan. En tal caso la señal será, evidentemente, un

polinomio trigonométrico, 
$$f(t) = \frac{a_0}{2} + \sum_{j=1}^N (a_j \cos \frac{2\pi j t}{T} + b_j \sin \frac{2\pi j t}{T})$$
. Ahora bien, cualquiera de estas funciones queda unívocamente determinada por su valor en  $2N+1$  instantes dentro del intervalo  $[0, T)$ , pues comparten las propiedades de

los polinomios algebraicos ordinarios. De hecho, como  $\cos \theta = \frac{e^{i\theta} + e^{-i\theta}}{2}$  y  $\sin \theta = \frac{e^{i\theta} - e^{-i\theta}}{2i}$ , para ciertos coeficientes complejos  $c_k$ , tenemos que

$$f(t) = \sum_{k=-N}^N c_k e^{\frac{2\pi i k t}{T}}$$
 . En particular, el polinomio algebraico 
$$P(z) = \sum_{k=0}^{2N} c_{k-N} z^k$$

satisface la identidad  $P(e^{\frac{2\pi i t}{T}}) = e^{\frac{2\pi i N t}{T}} f(t) =: g(t)$ . Se sigue que tanto  $g(t)$  como

$f(t) = e^{\frac{-2\pi i N t}{T}} g(t)$  quedan completamente determinadas a partir de cualquier

conjunto de  $2N+1$  valores  $f(t_k)$  con  $t_k \in [0, T)$ , pues dichos valores determinan unívocamente el polinomio  $P(z)$ . De hecho, ya en 1765, J. L. Lagrange había obtenido una fórmula explícita para el polinomio  $P(z)$  en términos de los valores  $g(t_k)$ , y el caso que acabamos de mostrar fue también estudiado, en 1841, por A.

Cauchy.

Si imponemos que los valores  $t_k$  se correspondan con un proceso de muestreo uniforme, tendremos que  $t_k = kh$  para cierto valor  $h > 0$ . Además, la condición de que  $\hat{x}(\xi) = 0$  para  $|\xi| > W$  se traduce, en este contexto, en que el grado  $N$  del polinomio trigonométrico que define a  $f$  verifica  $N/T \leq W$ . Por tanto,  $t_{2N+1} =$

$(2N+1)h < T$  implica que  $h < \frac{T}{2N+1} < \frac{T}{2N} < \frac{1}{2W}$ , que es lo que buscábamos. De

hecho, cualquier  $h < \frac{1}{2W}$  define al menos  $2N+1$  puntos distintos en  $[0, T)$  donde conocemos el valor de  $f(t_k)$  (aunque podrían no estar equiespaciados pues, si  $(2N+1)h < T$  entonces  $t^* = (2N+1)h - T$  puede diferir de los otros valores  $t_k$  donde conocemos  $f$ ), por lo que siempre podemos acudir al polinomio de interpolación de Lagrange para concluir la demostración del caso periódico.

Si la señal de partida no es periódica, pero su transformada de Fourier se anula fuera de cierto intervalo  $[-W, W]$ , entonces también es posible recuperarla a partir de un conjunto discreto de muestras, las cuales pueden tomarse, además, equiespaciadas, aunque esta vez se requieren infinitos valores de la señal original. Este es el contenido del famoso teorema del muestreo.

La prueba se basa en una combinación de ideas que involucran, de forma muy cuidadosa, el teorema integral de Fourier y el desarrollo en serie de Fourier de una función periódica. Concretamente se procede del siguiente modo: por hipótesis, la señal  $f$  es de banda limitada, lo que significa que,  $\mathcal{F}(f)(u) = 0$  para todo valor  $u$  que esté fuera de cierto intervalo  $[-W, W]$ . El menor  $W \geq 0$  con esta propiedad se denomina ancho de banda de la señal. Se sigue que, si  $F$  denota la extensión  $2W$ -periódica de  $\mathcal{F}(f)|_{[-W, W]}(u)$ , la restricción de  $\mathcal{F}(f)$  al intervalo  $[-W, W]$ , entonces  $F$  es una función periódica de energía finita y, por tanto, coincide con su desarrollo en serie de Fourier,

$$F(u) = \sum_{k=-\infty}^{\infty} c_k(F) e^{\frac{2\pi i k u}{2W}} = \sum_{k=-\infty}^{\infty} c_k(F) e^{\frac{\pi i k u}{W}}.$$

Ahora bien, los coeficientes de Fourier de  $F$  están dados por

$$\begin{aligned} c_k(F) &= \frac{1}{2W} \int_{-W}^W F(u) e^{-\frac{2\pi i k u}{2W}} du = \frac{1}{2W} \int_{-W}^W \mathcal{F}(f)(u) e^{-\frac{2\pi i k u}{2W}} du \\ &= \mathcal{F}^{-1}(\mathcal{F}(f))\left(\frac{-k}{2W}\right) = f\left(\frac{-k}{2W}\right). \end{aligned}$$

Se sigue que la transformada de Fourier de  $f$  queda completamente determinada a partir de los valores que toma la señal en los puntos  $k/(2W)$  con  $k \in \mathbb{Z}$  y, por tanto, si tenemos en cuenta otra vez el teorema integral de Fourier, también la propia señal  $f$  queda determinada a partir de estos valores. ¡Y eso es esencialmente lo que buscábamos!

La primera vez que se enunció un teorema de este tipo fue en un artículo del matemático inglés E. T. Whittaker, en 1915. Él se interesó por determinar una función  $C(t)$  que verificase las condiciones de interpolación  $C(a+kT)=f_k$  para todo  $k \in \mathbb{Z}$  y cierta sucesión de valores  $f_k$ . De hecho, propuso como solución utilizar la serie

$$C(t) = \sum_{k=-\infty}^{\infty} f_k \frac{\sin\left(\frac{\pi}{T}(t-a-kT)\right)}{\frac{\pi}{T}(t-a-kT)}.$$

Solo cinco años después, en 1920, K. Ogura demostraría que si  $f(t)$  es una señal de energía finita cuya transformada de Fourier se anula fuera de  $\left[-\frac{b}{2\pi}, \frac{b}{2\pi}\right]$ , entonces esta se puede recuperar completamente a partir de sus valores en los puntos  $k\pi/b$  con  $k \in \mathbb{Z}$  gracias a la fórmula

$$f(t) = \sum_{k=-\infty}^{\infty} f(k\pi/b) \frac{\sin(bt - k\pi)}{bt - k\pi} = \sin(bt) \sum_{k=-\infty}^{\infty} f(k\pi/b) \frac{(-1)^k}{bt - k\pi}$$

$$T = \frac{\pi}{b} = \frac{1}{2} \frac{1}{b/(2\pi)}$$

que se obtiene de la fórmula de Whittaker imponiendo  $a = 0$  y .  
Resumiendo, si  $W=b/(2\pi)$  es el ancho de banda de la señal, entonces esta se

puede recuperar a partir de muestras equiespaciadas cuando el muestreo se realiza con la frecuencia  $2W$ , que se llama frecuencia de Nyquist, en honor al físico e ingeniero H. Nyquist, quien formuló el mismo principio en forma de conjetura en un artículo de 1928.

En este punto es crucial advertir que el teorema funcionará del mismo modo si asumimos un ancho de banda  $\tilde{W}$  superior al verdadero ancho de banda de la señal (pues obviamente si  $\tilde{W} > W$  entonces  $[-W, W] \subseteq [-\tilde{W}, \tilde{W}]$  y, por tanto, la transformada de Fourier de  $f$  se anula también fuera de  $[-\tilde{W}, \tilde{W}]$ ). Se sigue que la fórmula anterior, en la que recuperamos la señal a partir de sus muestras, se puede aplicar con toda tranquilidad si sobremuestreamos, es decir, si la frecuencia de muestreo es superior a la frecuencia de Nyquist. Es más, se puede probar que el uso de una frecuencia de muestreo superior a la mínima exigible tiene la ventaja de reducir el efecto de posibles errores de redondeo, es decir, que produce un teorema más estable que el teorema original. Por otra parte, como veremos, si intentamos recuperar la señal con una frecuencia de muestreo inferior a la frecuencia de Nyquist, no tendremos éxito.

Este resultado fue demostrado independientemente por V. K. Kotelnikov en 1933 —por lo que en Rusia se le atribuyó durante mucho tiempo— y, en 1948 fue introducido en el contexto de la ingeniería por C. Shannon, quien lo usó en sus trabajos fundacionales sobre teoría de la información. Actualmente se conoce como teorema del muestreo de Shannon-Whittaker-Kotelnikov. En realidad, la lista de nombres que se le pueden asociar, con toda justicia, es mucho más larga y, de hecho, existen varios artículos dedicados a descifrar en todos sus detalles la historia del descubrimiento de este teorema. Además, existen muchas variaciones del resultado. Por ejemplo, se puede probar que la misma fórmula que hemos expuesto aquí se satisface con distribuciones temperadas cuya transformada de Fourier tiene soporte compacto. Es más, si pretendemos que el teorema sea útil para aplicaciones en física e ingeniería, esta versión distribucional es imprescindible, pues las distribuciones temperadas

representan un caso lo bastante general como para que contenga como casos particulares aquellas situaciones en las que el teorema requiere ser usado en la práctica. En efecto: la fórmula que describe la síntesis de  $f$  a partir de sus muestras, tal como la hemos demostrado hace un momento, requiere que la función  $f$  sea de energía finita, cosa que no sucede con las funciones  $\sin 2\pi kt$ ,  $\cos 2\pi kt$ , que tanta importancia tienen en las aplicaciones. Sin embargo, estas funciones, vistas como funciones generalizadas, son distribuciones temperadas con transformada de Fourier de soporte compacto y, por tanto, a ellas se les puede aplicar el teorema del muestreo en su versión distribucional. Otro caso al que podemos aplicar el mismo argumento, incluso sin aludir a la fórmula del muestreo, ya que aún no habíamos probado que estas funciones se pueden recuperar a partir de un conjunto discreto de muestras, es el caso de las funciones casi periódicas de banda limitada.

Finalmente, si eliminamos la exigencia de tomar muestras uniformemente, también hay una larga lista de resultados a los que podemos acudir. Además, estos teoremas se pueden demostrar para señales que dependen de un número finito de variables, y también existen resultados de este tipo en el mundo digital.

Quizás vale la pena citar aquí unas frases de Lüke, quien en un artículo de 1999 comentaba: “[La historia del teorema del muestreo] revela un proceso que aparece con frecuencia asociado a los problemas teóricos de la física y la tecnología: primero, los profesionales introducen una ‘regla de oro’ [en su trabajo práctico], entonces los teóricos desarrollan una solución general del problema en cuestión y, finalmente, alguien descubre que los matemáticos habían resuelto el problema matemático que hay oculto tras dicha propiedad hace tiempo, pero en condiciones de completo aislamiento”.

## El gran milagro: uso de la transformada rápida de Fourier

Una vez hemos asegurado que tomando muestras con suficiente frecuencia no se produce pérdida de información, y si tenemos en cuenta, además, que se conocen perfectamente las frecuencias máximas a las que trabaja el encéfalo, resulta

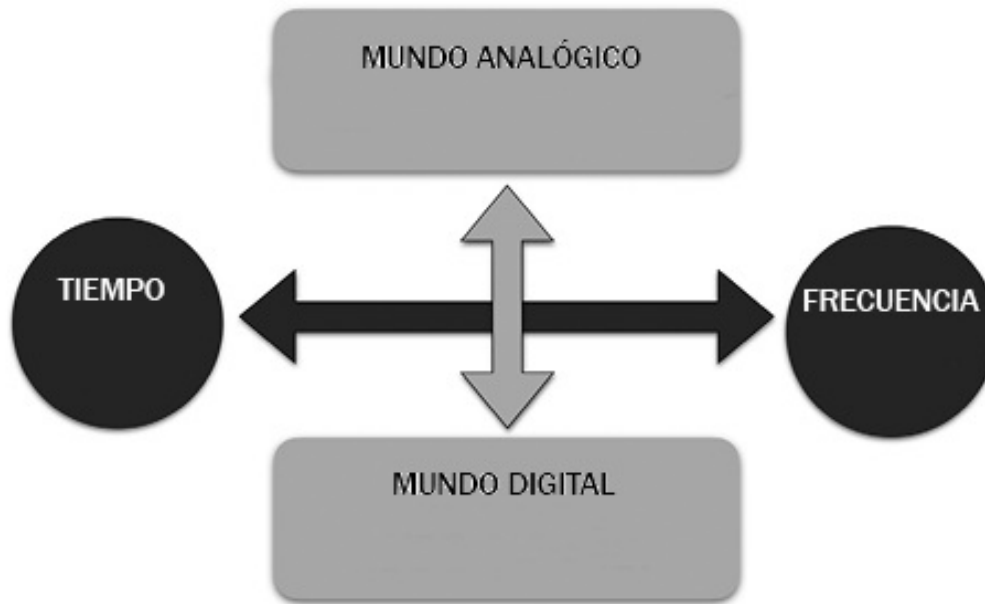
evidente que los sistemas de grabación de EEG solo requieren tomar un número limitado de muestras por segundo de la actividad eléctrica. Es decir, cada electrodo de cualquier sistema de grabación de EEG guarda su información en forma de un único vector de muestras, el cual tendrá con toda seguridad un tamaño importante, pero no deja de ser un vector. En este contexto cabe, pues, realizar el análisis de Fourier propio de los vectores, análisis que se resume, en esencia, en el cálculo de la dft de dicho vector. Obviamente, por motivos de tamaño de los vectores, los algoritmos que se aplican constantemente en este contexto son las “transformadas rápidas de Fourier”, que aceleran los cálculos en gran medida, pero el objeto que calculan es la dft, que tiene interés y significado teóricos por sí mismo. Nosotros no vamos a enredarnos en análisis de eficiencia de los cálculos, sino que preferimos enfatizar su significado físico.

Sin embargo, hay que tener en cuenta que todo nuestro interés se centra en el estudio de los EEG y que en relación con ellos los análisis de frecuencias que resultan relevantes no son aquellos que de manera natural aparecen en  $\mathbb{C}^N$  sino que lo son en términos de las frecuencias analógicas de las señales eléctricas que hemos muestreado. La pregunta es, por tanto, ¿hasta qué punto el uso de la dft de los datos muestreados nos proporciona información fidedigna del comportamiento frecuencial de estas señales cuando se consideran como señales analógicas? Porque, además, los algoritmos en el mundo digital son, en comparación con los cálculos correspondientes en el mundo analógico, mucho más simples y, por tanto, no podemos renunciar de ningún modo a hacer estos cálculos con la dft. De hecho, la mayor parte de las veces resulta imposible realizar los cálculos con funciones ordinarias en el mundo analógico y, sin embargo, ¡la dft siempre se puede calcular! Que la dft se pueda utilizar para estos fines es, ciertamente, un hecho maravilloso y afortunado, del que la comunidad científica debería alegrarse y regocijarse cada día. ¡Es la mayor de las suertes y nosotros vamos a aprovecharnos de ello! En los siguientes párrafos vamos a esbozar algunas de las razones que explican por qué (y cómo) podemos

usar la dft en las aplicaciones.

FIGURA 6

El teorema del muestreo permite pasar del mundo analógico al digital. La transformada de Fourier admite una versión en ambos mundos, para pasar del dominio del tiempo al dominio de las frecuencias. Los cálculos siempre se hacen en el mundo digital. Las frecuencias de interés son, sin embargo, siempre las analógicas.



Para empezar, debe quedar claro que si no tocamos para nada la dft del vector de muestras, entonces esta, por sí misma, mostrará un aspecto raro cuando se intenta interpretar en términos de frecuencias analógicas. Será un objeto extraño y, hasta cierto punto, ininteligible. No tendrá (aparentemente) nada que ver con lo que sería una interpretación natural de la señal analógica, tal como la conocemos. Nos dará la impresión de que el “análisis de Fourier en dimensión finita” es algo completamente diferente del “análisis de Fourier clásico” (que sí se corresponde estrechamente con conceptos de la física y sí sabemos cómo interpretar) y, por tanto, inútil en cuanto a su uso para aplicaciones “reales”. Afortunadamente, esta impresión inicial es solo apariencia. Ahora bien, es necesario familiarizarse con un conjunto de técnicas y herramientas que nos van a permitir realizar el paso deseado, que no es otro sino interpretar



adecuadamente, en términos analógicos, los cálculos que se realizan en el mundo digital.

Veamos algunos de los aspectos que debemos tener en cuenta. Por simplicidad, restringimos nuestra atención al caso unidimensional, pero este tipo de problemas, así como las técnicas que ayudan a resolverlos, se dan también en el caso multidimensional. Los problemas a los que debemos enfrentarnos son, en líneas generales, los siguientes:

- Reorganización de las salidas y de las entradas
- Reescalado de los ejes para entradas y salidas
- Problemas relacionados con el muestreo (*aliasing*) y con la finitud (*leakage*) de la señal.

Al calcular la dft de un vector, típicamente consideramos que este proviene de tomar muestras de una cierta señal analógica  $x(t)$  (y, de hecho, tal es el caso con los datos recogidos por cada electrodo en la grabación de los EEG) Es decir, nuestro vector es de la forma  $x=(x_0, x_1, \dots, x_{N-1})$  con  $x_k=x(kT)$  para cierto periodo de muestreo  $T>0$  y para  $k=0,1,\dots,N-1$ . Entonces, la dft del vector  $x$  contiene la información en frecuencias del mismo, desde el punto de vista del análisis de Fourier en dimensión finita. Lo interesante es que dicho vector ¡también contiene la información en frecuencias analógicas de la señal original! Sin embargo, la identificación de estas frecuencias no es trivial. Una primera razón es que ellas no aparecen ordenadas del modo convencional. Si  $y=dft(x)=(y_0, \dots, y_{N-1})$ , resulta que  $y_0$  representa la frecuencia cero, que es la componente de corriente continua de la señal, si esta se interpreta en términos de transmisión eléctrica. Las frecuencias  $y_k$ , para  $k=1,2,\dots,N/2-1$  son, en efecto, las componentes frecuenciales primera (que denotamos por  $f_1$  y que se corresponde con la frecuencia fundamental de la señal, es decir, la frecuencia más baja que interviene en la descripción de la señal, y que tiene la propiedad de que el resto de frecuencias son múltiplos suyos), segunda (es decir,  $2f_1$ ), tercera (esto es,  $3f_1$ ),

etc. Pero a partir del valor  $N/2$ , las frecuencias aparecen invertidas y se corresponden con componentes frecuenciales negativas. Concretamente,  $y_{N-1}$  se corresponde con la frecuencia  $-f_1$  y, en general,  $y_{N-k}$  se corresponde con la frecuencia  $-kf_1$  para  $k=1,2,\dots,N/2$ . Aquí estamos suponiendo que  $N$  es par, lo cual no implica una restricción severa en ningún caso de interés.

Si tenemos lo anterior en consideración, resulta natural que, a la hora de representar gráficamente la dft, es necesario que los datos se reorganicen intercambiando la primera mitad del intervalo con su segunda mitad simplemente deslizando la primera mitad hacia la derecha y la segunda mitad hacia la izquierda. El centro del nuevo intervalo se corresponderá con la frecuencia cero, las frecuencias negativas estarán a su izquierda y las positivas a su derecha.

A veces es también necesario reorganizar los datos contenidos en el propio vector de entrada, al cual se le va a calcular la dft. Esto es así también por motivos relacionados con la geometría de la definición de la dft. Cuando se define la dft del vector  $x$  se supone que  $x_k=x(kT)$  para  $k=0,1,\dots,N-1$ . Por tanto, si se han tomado las muestras de otra forma (por ejemplo, si hemos tomado muestras en un intervalo simétrico respecto del origen), entonces será necesario reorganizar los datos contenidos en el vector  $x$  para que estos verifiquen la ecuación anterior. En estos casos es, pues, deseable realizar una reorganización previa del vector  $x$  antes de calcular su dft, si queremos que luego esta represente de forma aproximada el comportamiento en frecuencias de la señal analógica original.

Además, al realizar el escalado de los ejes, debemos tener en cuenta la forma en la que se ha hecho el muestreo.

Por ejemplo, no pintamos el valor  $x_k$  sobre la abscisa de valor  $k$  sino sobre la abscisa cuyo valor es  $kT$ . Y, en el caso de la dft, debemos conocer el valor de la frecuencia fundamental  $f_1$ . Hay varias formas de deducir dicho valor. Veamos una: Si hemos tomado muestras con periodo de muestreo  $T$ , entonces la

frecuencia de muestreo usada es  $f_s=1/T$ , la cual, por el teorema del muestreo, debe ser al menos el doble que la frecuencia máxima que interviene en la señal. Por tanto, dando la vuelta al argumento, se concluye que con esta frecuencia de muestreo no podremos obtener información sobre las posibles frecuencias que intervienen en la señal para frecuencias superiores a  $1/(2T)$  y que, si la señal analógica original contiene estas frecuencias, entonces los datos proporcionados por su versión digital solo serán útiles para recuperar un alias de la señal original. De hecho, la dft nos proporciona información sobre las frecuencias entre  $-1/(2T)$  y  $1/(2T)$  y, como solo disponemos de  $N$  valores, los cuales representan muestras equiespaciadas de la transformada de Fourier analógica, resulta que la frecuencia fundamental debemos situarla en  $f_1=1/(NT)$ . Por tanto, los valores de la dft se deben representar entre los puntos  $-(N/2)f_1=-1/(2T)$  y  $(N/2)f_1=1/(2T)$ . Además, los valores  $y_k$  de la dft deben escalarse multiplicando por el periodo de muestreo  $T$ .

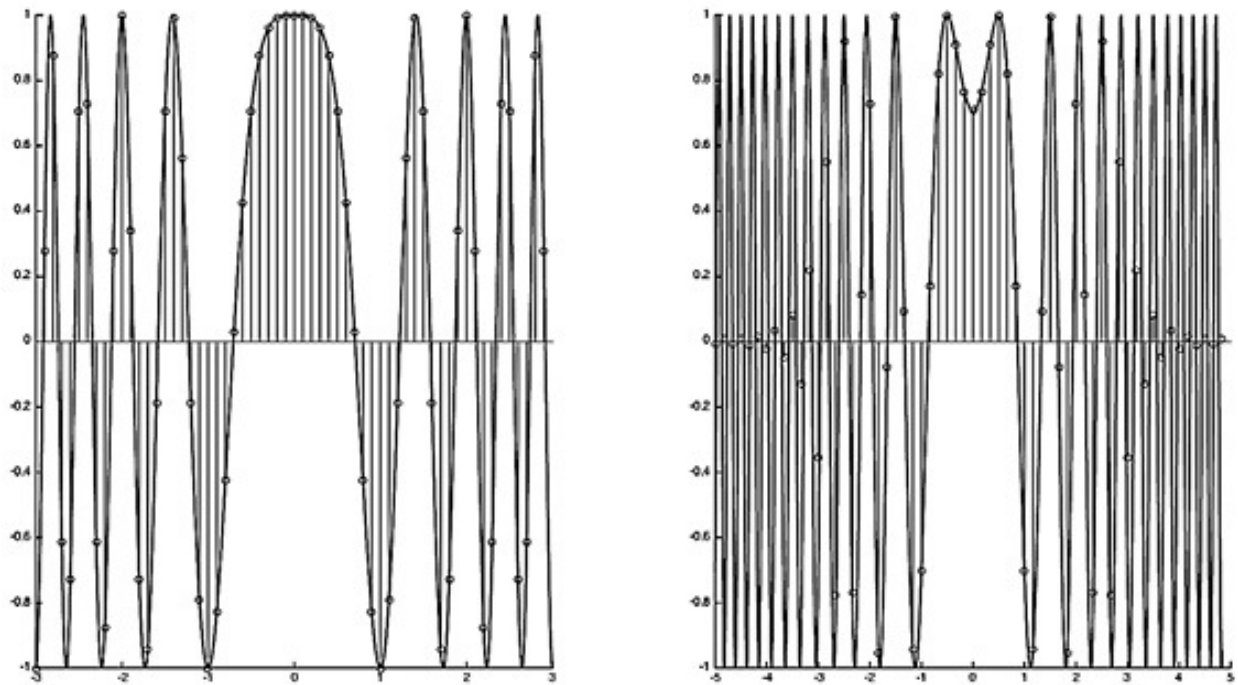
A pesar de que todos los cambios propuestos anteriormente proporcionan, en efecto, un valor de la dft que aproxima bastante bien al valor real de la transformada de Fourier analógica de la señal original, existen sin embargo algunas limitaciones en los cálculos, vinculadas a la propia naturaleza de la dft. En particular, al tomar muestras con cierta frecuencia de muestreo fijada de antemano, como hemos advertido insistentemente, se limita el espectro de frecuencias que podemos describir. Si la señal no es de banda limitada o si lo es, pero tomamos una frecuencia de muestreo inferior a la frecuencia de Nyquist, entonces no podremos recuperar la señal original sino solo una aproximación de la misma. Este fenómeno recibe el nombre de aliasing. Por otra parte, aún en el caso de que nuestra señal original sea de banda limitada, y que hayamos muestreado con la frecuencia exigida por el teorema del muestreo analógico, este nos pide un conjunto infinito de muestras equiespaciadas, algo que ciertamente nunca podremos satisfacer. Por tanto, la naturaleza finita de los cálculos que se pueden realizar implica la aparición de un cierto desajuste en los resultados, que

recibe el nombre de derramamiento espectral, manchado espectral, o leakage y es imposible evitar.

FIGURA 7

$x(t) = \cos(\pi t^2)$  muestreada a 8 Hz y cálculo de su transformada de Fourier

$\hat{x}(u) = \cos(\pi u^2 - \frac{\pi}{4})$  a partir de la dft.



El fenómeno de aliasing se puede controlar esencialmente de dos modos: bien incrementamos la frecuencia de muestreo, acercándonos lo máximo posible a la frecuencia de Nyquist, o bien filtramos la señal original con un filtro paso bajo que se adapte a nuestra velocidad de muestreo, y muestreamos la señal filtrada. Obviamente, ambas técnicas se pueden combinar. Aun así, es importante recordar que, en la grabación de los EEG, se toman siempre las señales con sobremuestreo. De hecho, lo habitual es muestrear al menos a 250 Hz, y se está imponiendo hacerlo a 500 Hz o 1.000 Hz. En algunos casos se usan frecuencias de hasta 2.000 Hz. Existen poderosas razones para realizar el sobremuestreo. A nosotros nos interesa sobremuestrear porque de esta forma se aumenta el cociente señal-ruido SNR, que es una medida estándar de la calidad de la señal y

se calcula dividiendo la energía de la señal limpia por la energía del ruido introducido. Obviamente, a mayor cociente SNR, mejor calidad de la señal. El sobremuestreo también permite una mejor estimación de la actividad en altas frecuencias. En particular, si queremos tener seguridad de que la dft del vector de muestras sea representativa del comportamiento en frecuencias de la señal original, el sobremuestreo es, excepto en contadas excepciones, imprescindible.

El derramamiento espectral se debe a que, al imponer un número finito de muestras, sustituimos la señal a la que se podría aplicar el teorema del muestreo por la que resulta de multiplicar la señal original por una función del tipo  $w(t)=\text{rect}(t/R)$ , donde  $\text{rect}(t)=1$  para  $0 \leq t \leq 1$  y  $\text{rect}(t)=0$  en el resto de casos. Esto puede producir, por tanto, la aparición de una discontinuidad de salto entre el principio y el final de la señal que vamos a muestrear, provocando que aparezcan, en el dominio de la frecuencia, componentes de altas frecuencias que antes no estaban presentes. Es más, si tenemos en cuenta el principio de incertidumbre (que exponemos más adelante), incluso en el caso de que no aparezcan discontinuidades, la nueva señal que estamos muestreando no puede ser de banda limitada, pues está limitada en el dominio del tiempo. De hecho, en el dominio de la frecuencia, el efecto del truncamiento es, gracias a la descripción anterior y al teorema de convolución, el resultado de convolucionar la transformada de Fourier de la señal original con la transformada de Fourier de  $w(t)$ , que es la señal  $W(u)=\sin(\pi Ru)/(\pi u)$ . Esta convolución puede causar problemas de dos tipos, esencialmente. Puede suavizar y extender la señal,

debido al lóbulo principal de la señal seno cardinal (es decir, la función  $\frac{\sin(\pi u)}{\pi u}$ ), y también puede que, como consecuencia de la convolución con los lóbulos menores de la función seno cardinal, se formen oscilaciones en altas frecuencias, que irán decayendo con el paso del tiempo (este fenómeno se llama rizado). Si la transformada de Fourier de la señal original es suave y no contiene impulsos, entonces el derramamiento espectral será mínimo y, en muchos casos, resultará casi imperceptible. Sin embargo, si esta es una señal que contiene impulsos (en

otras palabras, si la señal original es periódica o casi periódica), entonces los efectos del derramamiento espectral se harán patentes.

Existen dos situaciones en las que el derramamiento espectral no se produce. En primer lugar, si la señal  $x(t)$  es de duración finita y su soporte está contenido en el soporte de la función  $w(t)=\text{rect}(t/R)$ , entonces al multiplicar ambas funciones no se produce cambio alguno. El otro caso que no da problemas es cuando  $x(t)$  es periódica y la longitud  $R$  de la ventana definida por  $w(t)=\text{rect}(t/R)$  es un múltiplo entero del periodo de  $x(t)$ . En el resto de casos, el derramamiento espectral es inevitable, y lo único que podemos hacer es reducir su efecto lo máximo posible.

Para lograr esto existen dos técnicas que, como en el caso del aliasing, pueden combinarse entre sí. La primera es aumentar el tamaño de  $R$ , lo cual nos obliga a tomar muestras de la señal en un rango más amplio, y tiene como consecuencia que el primer lóbulo de la función seno cardinal se estrecha —esto provoca que la convolución contra ella se parezca cada vez más a la convolución contra un impulso, que nos recupera la señal original— y los siguientes lóbulos disminuyen de tamaño muy rápidamente —lo cual tiene el efecto de disminuir drásticamente el tamaño de las oscilaciones en altas frecuencias—. Otra técnica que se suele utilizar es el llamado estrechamiento espectral, que consiste en multiplicar la función original  $x(t)$  por una función (que llamamos “ventana”)  $v(t)$  cuyo soporte esté contenido en  $[0,R]$ , que tome valores cercanos a 1 en la mayor parte del soporte, y luego decaiga con la mayor suavidad posible. En otras palabras, se trata de sustituir la ventana  $w(t)=\text{rect}(t/R)$  por otra más suave (sin discontinuidades) que tenga el mismo soporte. El precio a pagar es que algunas de las muestras que se toman cerca del borde del intervalo  $[0,R]$  se verán perturbadas, pero la ganancia en el dominio de la frecuencia puede ser enorme, haciendo que casi desaparezcan las ondulaciones en altas frecuencias y, aun así, se respete bastante bien el comportamiento espectral de la señal en bajas frecuencias. De todas formas, se paga el precio de perder resolución espectral, es decir, se pierde cierta capacidad de distinguir el comportamiento en frecuencias

cuando estas están excesivamente cerca unas de otras. Existen numerosas ventanas que se utilizan para el estrechamiento espectral, dependiendo del contexto en el que se trabaje, y existe también bastante literatura dedicada a describir las ventajas e inconvenientes de cada una de ellas. Estas ventanas suelen estar implementadas en el *software* científico que se utiliza en las distintas aplicaciones existentes.

A partir de ahora vamos a interpretar, cuando hablamos de la dft de un vector de muestras  $x$  (tomado por ejemplo de las mediciones en un electrodo del EEG), que hemos calculado la información espectral analógica de la señal original siguiendo los criterios introducidos en esta sección.

## **Incertidumbres. Análisis en tiempo-frecuencia**

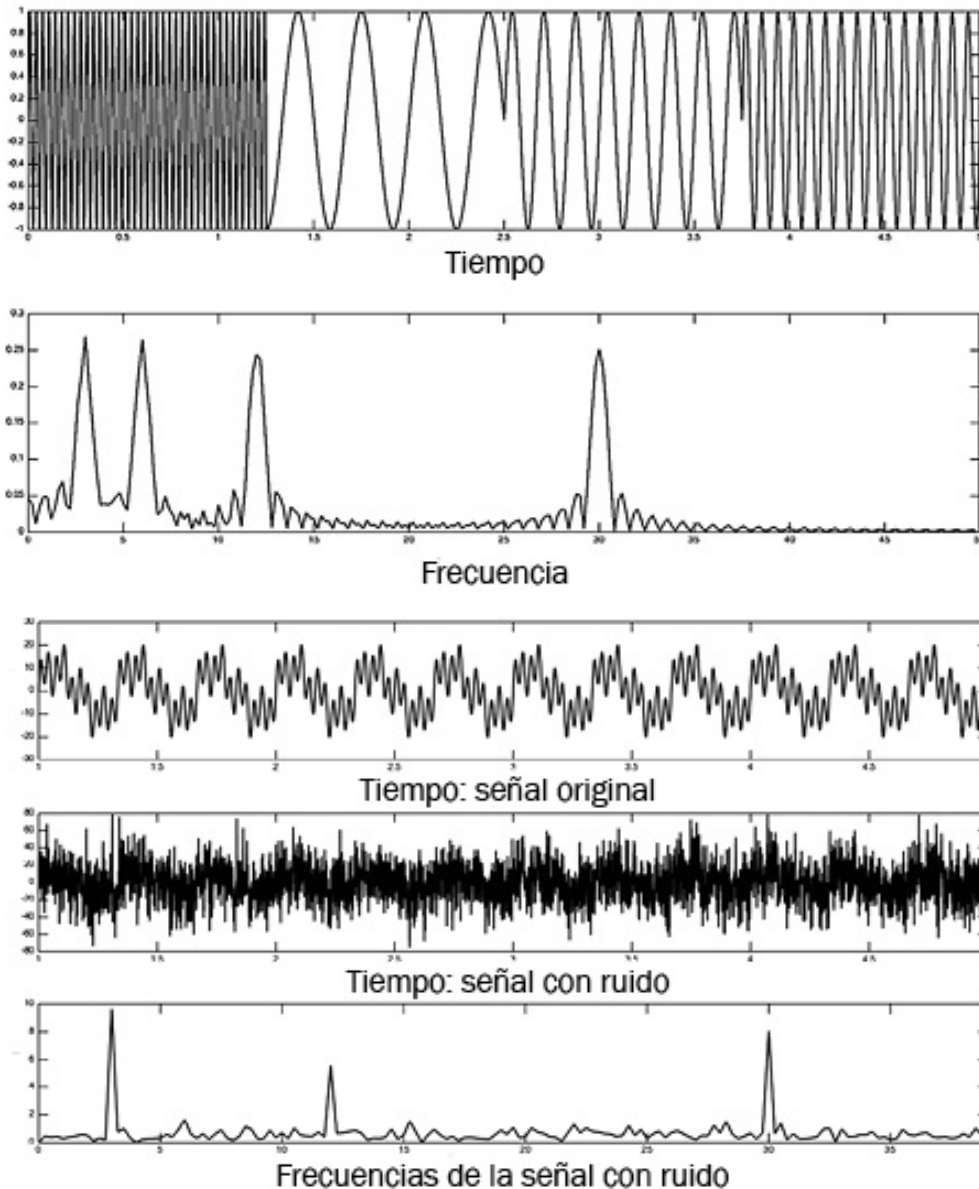
Las señales recogidas por cada uno de los electrodos de un electroencefalograma no encajan a la perfección con ninguno de los modelos clásicos de aplicación del análisis de Fourier, pues no son periódicas ni tampoco decaen en el infinito. De hecho, ni siquiera tiene sentido pensar en ellas como señales que se extienden infinitamente en el tiempo. Sin embargo, sí es cierto que presentan oscilaciones rítmicas, y, aunque dicho ritmo puede variar con el tiempo, este resulta relevante desde un punto de vista fisiológico y biomédico. Podríamos decir que el análisis de Fourier se puede aplicar a señales estacionarias, que son aquellas cuyo contenido en frecuencias no varía con el paso del tiempo y, en particular, cada una de sus componentes frecuenciales está presente en todo instante. Las señales cuyo contenido en frecuencias depende del tiempo se llaman no estacionarias. Aun así, resulta tentador aplicar el análisis de Fourier para el estudio de señales no estacionarias. ¿Cuándo es esto posible y de qué modo se puede proceder?

Vamos a mostrar lo que sucede al aplicar la dft al vector de muestras de una señal no estacionaria para algunos casos-tipo, que se corresponden con situaciones que, en efecto, podemos encontrarnos en el análisis del EEG. El primer caso que consideramos es una señal cuyo comportamiento en frecuencias varía con el tiempo: la señal oscila con frecuencias diferentes, bien definidas, en

diferentes intervalos temporales, también bien definidos. Consideramos también el caso de una señal periódica que resulta de la superposición de varios armónicos puros, a la que se le ha sumado ruido aleatorio.

FIGURA 8

Varios ejemplos de señales muestradas y las amplitudes de sus dft.

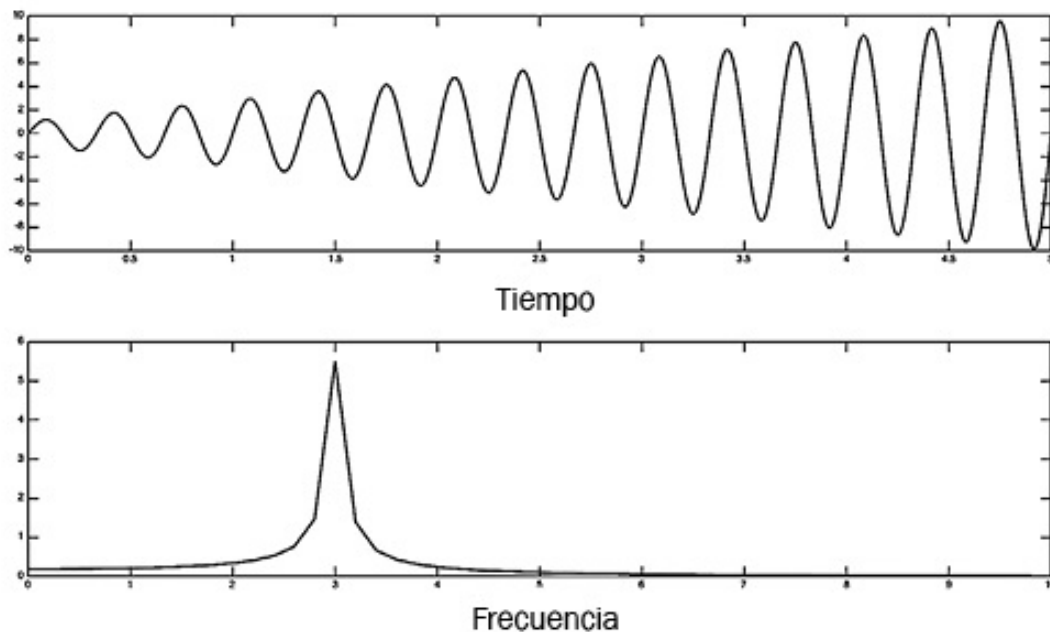


Como se ve, la dft capta en ambos casos cuáles son las frecuencias fundamentales que intervienen en la señal (en este caso, 3, 6, 12 y 30 Hz) pero



no proporciona información sobre su localización temporal y, además, aparecen otras frecuencias que, si no conocemos la señal original, podrían interpretarse como ruido o como altas frecuencias. El siguiente caso que consideramos es una señal cuya frecuencia de oscilación es la misma (3 Hz) a lo largo del tiempo, pero produce oscilaciones cuya amplitud va en aumento conforme pasa el tiempo (análogamente, podríamos considerar el caso en el que la amplitud de las frecuencias decrece con el transcurso del tiempo).

FIGURA 9  
Otro ejemplo de señal mostrada y las amplitudes de su dft.



Obviamente, la dft capta que la frecuencia principal es 3 Hz, pero incluye muchas otras frecuencias en torno a ella, y el aspecto del gráfico resultante no da excesivas pistas sobre el comportamiento en frecuencias de la señal original.

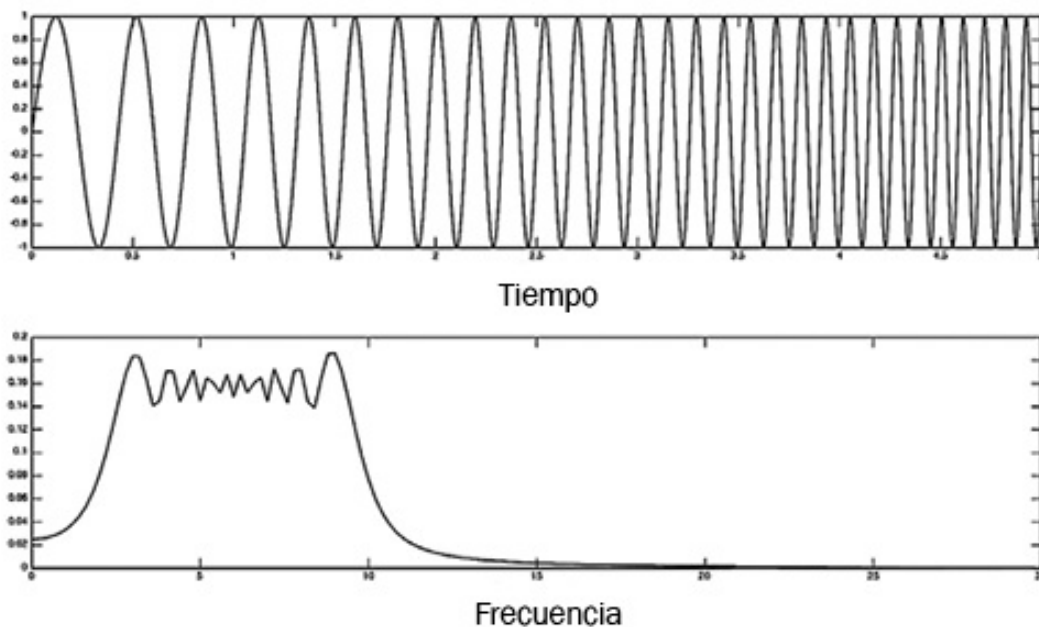
Otro caso es el de una señal que aumenta su frecuencia de vibración de forma progresiva, manteniéndose constante la amplitud máxima de las oscilaciones (es decir, no se produce un salto repentino en el cambio de las frecuencias sino que estas van en aumento progresivo).

Nuevamente, se observa que la dft no es especialmente útil para interpretar el comportamiento en frecuencias de este tipo de señales, pues su representación

gráfica no nos informa sobre el crecimiento gradual de las frecuencias (es decir, revela que existen frecuencias presentes, con similar amplitud, en un amplio rango de frecuencias, pero no nos dice absolutamente nada sobre su localización temporal (podrían estar todas ellas presentes todo el tiempo que dura la señal) y, además, aparecen otras frecuencias por encima y por debajo de las que realmente intervienen en la señal, y no queda claro si estas se pueden interpretar como ruido.

FIGURA 10

Otro ejemplo de señal mostrada y las amplitudes de su dft.



Existen, además, señales a las que no podemos aplicar el análisis de Fourier clásico a pesar de que son estacionarias. Un ejemplo de ello son las señales casi periódicas, es decir, aquellas que resultan de la superposición de varias señales periódicas con al menos dos de ellas con periodos cuyo cociente es un número irracional.

Afortunadamente, disponemos de técnicas apropiadas para el estudio de señales no estacionarias. De particular importancia son el análisis de Fourier enventanado y la teoría de ondículas. Las ideas que conducen a estas teorías son

relativamente sencillas. Multiplicamos la señal original  $f(t)$  por una función  $w(t-t_0)$ , que representa una “ventana centrada en el punto  $t_0$ ”, la cual se anula (o toma un valor verdaderamente despreciable, ínfimo) fuera de un cierto intervalo (conocido) centrado en  $t_0$ , y está cerca del valor 1 en la mayor parte de dicho intervalo. Así, cuando calculamos la transformada de Fourier del producto  $f(t)w(t-t_0)$  lo que hacemos es concentrar nuestra atención sobre el comportamiento en frecuencias de la señal  $f(t)$  en un pequeño entorno del punto  $t_0$ . Obviamente, se trata de la misma técnica que hemos descrito cuando hablamos de estrechamiento espectral. En otras palabras, dicha transformada nos dice qué frecuencias explican el comportamiento de la señal cerca de  $t_0$ . Se supone que esto lo podremos hacer si la señal bajo consideración es estacionaria en el trozo de intervalo donde la ventana no se anula, que debe ser lo bastante pequeño como para garantizar este hecho, pero lo bastante grande como para detectar las oscilaciones de  $f$ . Formalmente, estaríamos considerando la función de dos variables:

$$S\mathcal{F}(f)(t_0, \xi_0) = \mathcal{F}(f(t)w(t-t_0))(\xi_0) = \int_{-\infty}^{\infty} f(t)w(t-t_0)e^{-2\pi i \xi_0 t} dt,$$

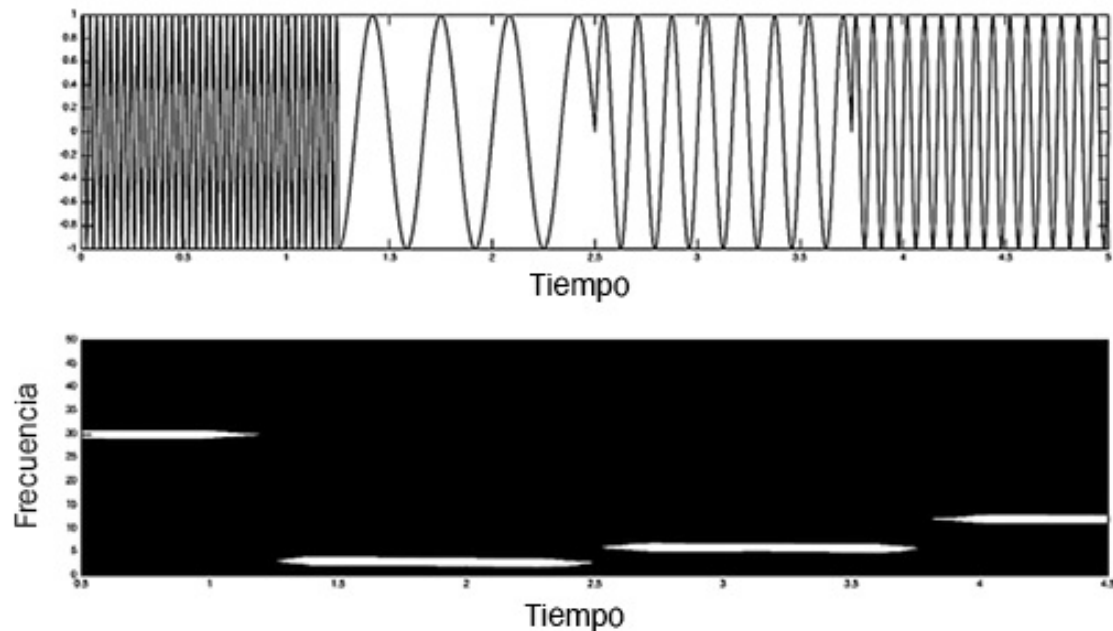
que recibe el nombre de transformada de Fourier enventanada de la señal  $f$ .

Supongamos que estamos interesados en conocer qué frecuencias intervienen en la descripción de nuestra señal en cada instante de tiempo. Con esta idea en mente, podemos dibujar en el plano tiempo-frecuencia el módulo de  $S\mathcal{F}(f)(t, \xi)$ , dando lugar a lo que se denomina un espectograma de  $f$ . A cada punto  $(t_0, \xi_0)$  del plano se le asigna un nivel de gris que refleja la amplitud de  $|S\mathcal{F}(f)(t_0, \xi_0)|$ . Asignamos al valor 0 el color negro y, a mayor amplitud, más claro pintamos el punto.

Por supuesto, al igual que sucedía con el análisis en frecuencia de las señales analógicas, estos cálculos no se pueden hacer, en general, por la sencilla razón de que la transformada de Fourier es un objeto muy difícil de calcular. Ahora bien,

la misma filosofía de trabajo que hemos seguido anteriormente, digitalizando nuestras señales y realizando un análisis de Fourier en el mundo digital, estos cálculos se pueden realizar con la dft del vector de muestras sin ninguna dificultad, y ellos son en efecto representativos de lo que sucede en el mundo analógico. Por ejemplo, si hacemos estos cálculos con la señal que consideramos anteriormente, la cual oscilaba con unas frecuencias bien definidas (de 3, 6, 12 y 30 Hz) en diferentes intervalos de tiempo, obtenemos un gráfico que, ahora sí, refleja con claridad el comportamiento en tiempo-frecuencia de la señal.

FIGURA 11  
Ejemplo de señal muestrada y su transformada de Fourier enventanada.



De todas formas, también se observa cierta dificultad, reflejada en el grosor de las líneas blancas que aparecen en la figura, así como en la aparente separación temporal entre unas líneas y otras, que nos impide ser precisos en la definición exacta de las frecuencias implicadas, así como en su localización temporal exacta. Esto se hace más nítido en la siguiente figura, donde se han hecho los

cálculos con la señal  $\sin(\frac{40\pi}{3}t^2)$ , cuyo comportamiento en frecuencias varía gradualmente con el tiempo, y hemos realizado el análisis en tiempo-frecuencia

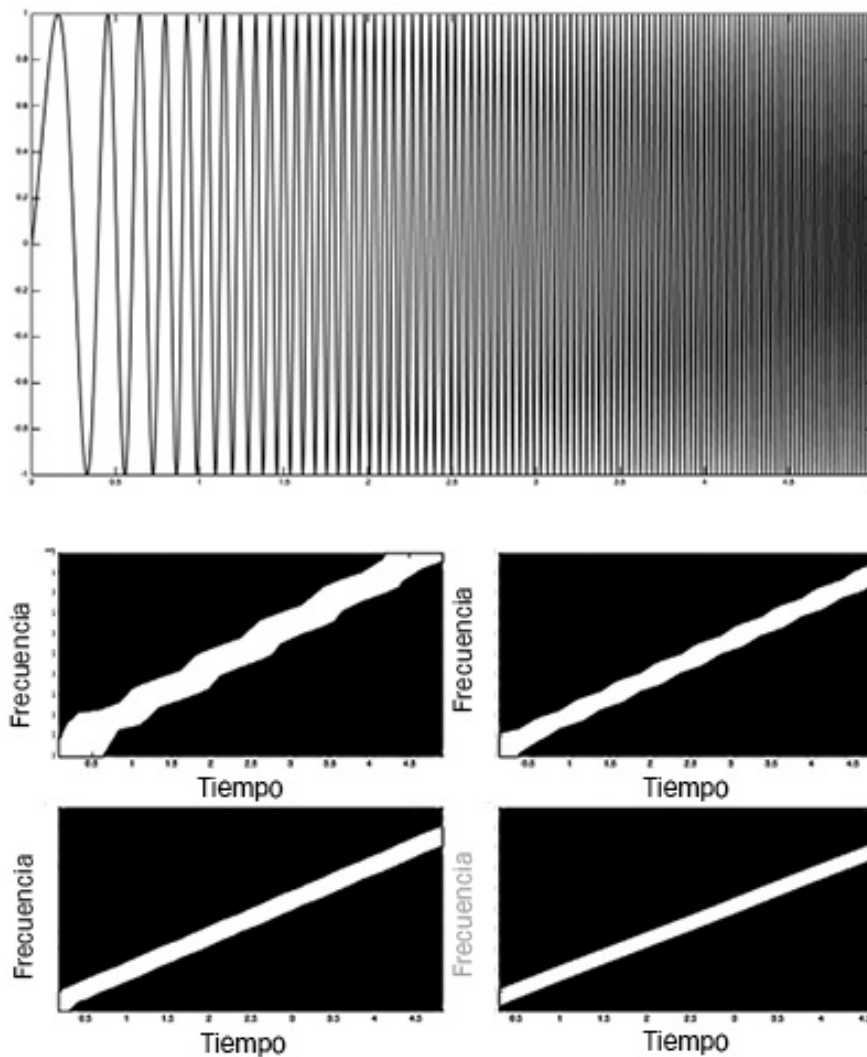
con la transformada de Fourier enventanada con una ventana (digital) de Hamming,

$$w(n) = \frac{1}{2} \left( 1 - \cos \frac{2\pi n}{N-1} \right),$$

para diferentes tamaños temporales  $N$  de la ventana. Concretamente, se toman los valores  $N=150, 250, 350, 600$  (de arriba abajo, de izquierda a derecha).

FIGURA 12

La precisión del análisis en tiempo-frecuencia depende del tamaño de la ventana.



La razón principal por la que el análisis en tiempo-frecuencia no se puede realizar sin introducir ambigüedades es el famoso principio de incertidumbre.

Este principio tiene especial relevancia en física de partículas e incluso en filosofía de la física. En matemáticas, sin embargo, se resume en establecer cierta desigualdad para la transformada de Fourier que, además, no es especialmente complicada. Concretamente, se trata de demostrar la desigualdad  $\Delta_{t_0} f \cdot \Delta_{\xi_0} \mathcal{F}(f) \geq \frac{1}{4}$  para toda señal de energía finita  $f$ , donde, si  $x(t)$  es una señal de energía finita,

$$\Delta_a x = \frac{\int_{-\infty}^{\infty} |(t-a)x(t)|^2 dt}{\int_{-\infty}^{\infty} |x(t)|^2 dt}$$

representa una medida de lo concentrada que está la señal  $x(t)$  en torno al punto  $a$ . En efecto, si  $x(t)$  es muy pequeño para los valores de  $t$  que se alejan un poco de

$a$ , entonces el factor  $(t-a)^2$ , hace que la integral  $\int_{-\infty}^{\infty} |(t-a)x(t)|^2 dt$  sea pequeña

en comparación a la integral  $\int_{-\infty}^{\infty} |x(t)|^2 dt$  y, por tanto,  $\Delta_a x$  será pequeño. Por otra parte, si la señal  $x$  está dispersa lejos de  $a$  (es decir, si contiene, para valores de  $t$  alejados de  $a$ , información relevante para su energía, que se computa con la

integral  $\int_{-\infty}^{\infty} |x(t)|^2 dt$ ), entonces el factor  $(t-a)^2$  hace que la integral  $\int_{-\infty}^{\infty} |(t-a)x(t)|^2 dt$  sea grande en comparación a la integral  $\int_{-\infty}^{\infty} |x(t)|^2 dt$  y, por

tanto,  $\Delta_a x$  será grande. La desigualdad  $\Delta_{t_0} f \cdot \Delta_{\xi_0} \mathcal{F}(f) \geq \frac{1}{4}$  significa, por tanto, que fijados un valor cualquiera  $t_0$  del tiempo y un valor cualquiera  $\xi_0$  de la frecuencia, es imposible que la señal  $f$  mantenga concentrada la mayor parte de su energía en torno a  $t_0$  y, simultáneamente, su transformada de Fourier  $\mathcal{F}(f)$  mantenga su energía concentrada en torno a la frecuencia  $\xi_0$ . En otras palabras: es imposible localizar con precisión infinita la señal simultáneamente en el dominio del tiempo y el dominio de la frecuencia. ¡No parece tan profundo, pero

sí lo es!

El efecto que tiene el Principio de Incertidumbre en el análisis de Fourier enventanado, es que la ventana  $w(t)$  determina el nivel de resolución dentro del cual podemos interpretar los resultados. Concretamente, los valores  $\Delta_{t_0} w$  y  $\Delta_{\xi_0} \mathcal{F}(w)$  determinan el tamaño de una caja, centrada en el punto  $(t_0, \xi_0)$ , con la propiedad de que la transformada enventanada  $S\mathcal{F}(f)$  es capaz de distinguir dos componentes frecuenciales de la señal que tengan presencia en el intervalo temporal  $t_0 \pm \Delta_{t_0} w / 2$  si están separadas por un valor superior a  $\Delta_{\xi_0} \mathcal{F}(w)$ , pero podría no distinguir las en otro caso. En resumidas cuentas, en la caja centrada en  $(t_0, \xi_0)$  de base  $\Delta_{t_0} w$  y altura  $\Delta_{\xi_0} \mathcal{F}(w)$  se resume la información en tiempo-frecuencia que podemos precisar para  $f$  en torno a  $(t_0, \xi_0)$ .

Aunque la transformada de Fourier enventanada implica un avance significativo respecto del análisis de Fourier clásico, aún persiste el problema de que los niveles de resolución fijados por cualquier ventana son los mismos en todo punto del plano tiempo-frecuencia, y esto podría no ser deseable para el estudio de algunas señales que contengan información en frecuencias en rangos muy diversos, y que esté presente en diferentes intervalos de tiempo. La teoría de ondículas resuelve este problema al sustituir el producto  $w(t - t_0)e^{-2\pi i \xi_0 t}$  que

aparece en la definición de  $S\mathcal{F}(f)(t_0, \xi_0)$  por una función (ondícula) del tipo

$$\varphi_{a,b}(t) = \frac{1}{\sqrt{a}} \varphi\left(\frac{t-b}{a}\right)$$

que determina niveles de resolución diferentes en función del valor  $a > 0$  ( $b \in \mathbb{R}$  nos informa sobre la localización temporal del análisis) y, simultáneamente, permite una cierta interpretación en términos de frecuencias. Concretamente, se asume que la función  $\phi$ , que llamamos “ondícula madre”, es una señal absolutamente integrable y de energía finita. Es más, se suele imponer

que  $\int_{-\infty}^{\infty} \varphi(t) dt = 0$  y  $\int_{-\infty}^{\infty} |\varphi(t)|^2 dt = 1$ . Esta función debe oscilar con una frecuencia baja (en torno a 1 o 2 Hz) y debe estar temporalmente concentrada en torno a  $t=0$  (de modo que lejos de este punto se anule o tenga un valor despreciable). Al considerar las versiones rescaladas y trasladadas  $\phi_{a,b}(t)$ , el valor de  $a$  nos informará sobre su contenido en frecuencias y el valor de  $b$  sobre su localización temporal.

La transformada correspondiente es, entonces,

$$\mathcal{WT}(f)(a,b) = \int_{-\infty}^{\infty} f(t) \overline{\phi_{a,b}(t)} dt$$

Hay muchas posibilidades para elegir ondículas. Un ejemplo especialmente importante para el estudio de los EEG es la familia de ondículas de Morlet, cuya ondícula madre, que depende de dos parámetros, está dada por

$$M_{c_0, \xi_0}(t) = e^{2\pi i \xi_0 t} e^{-\left(\frac{t}{c_0}\right)^2}$$

## Aplicaciones al EEG

Sabemos que es importante conocer los diferentes ritmos a los que oscila el EEG, pues se trata de una información que tiene significado neurofisiológico. También sabemos que la dft se puede utilizar para detectar dichos ritmos. Por tanto, sería conveniente explicar, muy brevemente, cómo se mide la actividad en cada una de las bandas de frecuencias, así como los algoritmos que permiten calcular algunos neuromarcadores, como la energía acumulada en una banda de frecuencias dada, o el cociente de Monastra  $R_{\theta/\beta}$ , que son frecuentemente utilizados en neuropsicología.

Si  $x$  representa el vector de muestras de uno de los electrodos del EEG, su dft nos informa sobre el contenido en frecuencias (analógicas) del potencial analógico original —del cual  $x$  se construye tomando muestras—. Por tanto, si nos interesa conocer la energía  $E_B$  acumulada en una banda de frecuencias



concreta  $B$ , lo único que tenemos que hacer es calcular la energía (es decir, el cuadrado de la norma Euclídea usual) del vector que resulta de tomar de la dft de  $x$  solamente los coeficientes correspondientes a las frecuencias que pertenecen a nuestra banda de frecuencias, haciendo 0 el resto. Este cálculo se puede realizar también del siguiente modo: se divide el intervalo temporal en épocas del mismo tamaño (por ejemplo, para el caso de mediciones realizadas sobre la actividad EEG espontánea una buena elección es imponer épocas de 2 segundos de duración, pero el tamaño de las épocas dependerá del estudio que estemos realizando). En cada época se calculan la dft y, tal como hemos dicho antes, la energía asociada a la banda de frecuencias elegida. Finalmente se promedian dichas energías, y ese es nuestro nuevo valor para la energía  $E_B$  de la señal en la banda  $B$ . Este modo de hallar el valor  $E_B$  es más estable (y más fiable) que el anterior y es, a todas luces, mejor que simplemente seleccionar un intervalo temporal concreto y hacer las cuentas para dicho intervalo (algo que algunos profesionales hacen, pero uno podría preguntarse entonces cuáles son los criterios para elegir un intervalo temporal frente a otro y cómo se pueden comparar los datos de los distintos estudios si esta elección es aleatoria. Aun así, obviamente, a veces es interesante seleccionar intervalos concretos porque ellos definen una situación especial, que influye en el EEG. Por ejemplo, no es lo mismo medir la actividad  $\alpha$  espontánea con los ojos cerrados que con los ojos abiertos, y podría ser importante medirla por separado en ambas circunstancias). Una vez ha quedado esto claro, el cálculo del cociente  $R_{\theta/\beta}$ , se hace simplemente calculando la energía en la banda theta,  $E_\theta$ , la energía en la banda beta,  $E_\beta$ , y dividiendo  $R_{\theta/\beta} = E_\theta/E_\beta$ .

En aplicaciones que requieran usar altas frecuencias, es conveniente realizar análisis en tiempo-frecuencia. Además, dicho análisis permite responder a preguntas relacionadas con la cantidad de tiempo en que cada electrodo vibra en cada banda de frecuencias, que es también un parámetro significativo para estudios neurocognitivos. Otro aspecto interesante es crear, para una frecuencia

determinada, un mapa topográfico de la cabeza que revele qué regiones han mostrado mayor (o menor) actividad en dicha frecuencia, durante el proceso de grabación del EEG (de modo que, por ejemplo, podamos determinar en qué zonas del encéfalo ha sido mayor la actividad  $\alpha$  y utilizar dicha información en nuestro estudio). Esto se hace calculando la dft de las señales registradas por todos los electrodos y, una vez seleccionada la frecuencia de interés, anotando sobre cada electrodo su contribución en dicha frecuencia concreta, información que luego se utiliza del modo habitual para dibujar el correspondiente mapa topográfico, el cual incluye el trazado de las curvas de nivel que definen dichas contribuciones sobre los distintos electrodos (véase el capítulo 3 para una explicación más detallada de los mapas topográficos, en otro contexto). También es frecuente, en muchos estudios, incluir la información espectral de los potenciales evocados o de sus componentes ICA (este concepto se introduce también en el próximo capítulo).

## CAPÍTULO 3

# Separación de fuentes y eliminación de artefactos

### Separación ciega de fuentes

Imagina que estás dando un agradable paseo por el campo, cerca de un río. Puedes escuchar el murmullo del agua. Pero también te llegan otros sonidos. Una moto que, a lo lejos, emite su estruendo de metal, un jilguero que canta hermosamente, tu acompañante que habla entre silencio y silencio. Tu cerebro es capaz de percibir estos distintos sonidos y, lo que es verdaderamente asombroso, los identifica y, si te concentras, puede aislarlos, separarlos. ¿Cómo hace esto? ¿Existe un mecanismo “automático” que separa las distintas fuentes sonoras? ¿Es el oído capaz de algo así?

Como sabemos, la propagación de ondas sonoras se rige, entre otras cosas, por el llamado principio de superposición. Cada sonido se puede representar con una señal que varía con el tiempo, y los distintos sonidos se superponen, se suman. Percibimos una suma de contribuciones que provienen de fuentes diferentes. El análisis de Fourier se basa precisamente en este principio acústico. Como hemos tenido ocasión de comentar, ya en tiempos de D. Bernoulli se esgrimió el argumento de la superposición de frecuencias distintas para garantizar que cualquier sonido y, de hecho, cualquier función periódica, se puede descomponer como suma (posiblemente infinita) de armónicos puros.

El problema propuesto aquí es, sin embargo, más complejo, puesto que no deseamos descomponer nuestra señal en armónicos puros sino que pretendemos separar las distintas fuentes sonoras tal como las percibimos. Queremos aislar el canto del jilguero y este con toda seguridad compartirá frecuencias con el sonido

del agua, incluso con la moto. Queremos separar sonidos que, para nosotros, tienen un significado preciso, y estos sonidos son complejos: involucran frecuencias en una amplia parte del espectro sonoro y estas frecuencias son utilizadas a la vez por las distintas fuentes sonoras que percibimos. Además, para complicar las cosas, las frecuencias varían con el tiempo. La conclusión inevitable es que el análisis de Fourier, por sí solo, no basta para resolver esta cuestión.

Si queremos responder al problema inicial, es decir, averiguar cómo se las ingenia nuestro cerebro para separar las distintas fuentes sonoras, lo razonable es argumentar que lo hemos entrenado, que responde a una experiencia acumulada. Al nacer (e, incluso antes, durante la gestación), somos capaces de percibir todo tipo de sonidos. Lo que no sabemos, lo que debemos aprender, es interpretarlos. Durante cierto tiempo, nos vemos obligados a prestar atención a los sonidos que percibimos para, con la experiencia, aprender a distinguir el silencio del sonido, el lenguaje de la música (u otros estímulos sonoros), lo que dice nuestra mamá de lo que dicen otras personas a nuestro alrededor, o el ladrido del perro, o la televisión. Aprendemos a distinguir si en un momento determinado se dirigen a nosotros o no, reconocer la cercanía del emisor de un sonido, el volumen del mismo, los diferentes tonos, etc. Toda esta información es acumulada y cuando somos sometidos a nuevas experiencias, adquirimos la capacidad de interpretarlas. No nacemos con el conocimiento intuitivo de lo que significan las distintas expresiones (acción, adjetivo, nombre, etc.) Algo parecido se produce con todos nuestros sentidos. Por ejemplo, es la experiencia acumulada la que nos permite interpretar, con solo la vista, el estado de ánimo de otras personas.

La cuestión de la experiencia se traduce en última instancia en términos estadísticos. Los ingenieros son capaces de diseñar máquinas que aprenden una conducta determinada mediante el uso de inteligencia artificial, y un punto clave para el éxito de este tipo de “inteligencia” (que no lo es obviamente, en sentido estricto), es el entrenamiento estadístico. La pregunta es, por tanto, la siguiente: ¿podemos diseñar algún tipo de algoritmo que responda de forma instantánea al

problema de la separación ciega de fuentes?

Si, como en el caso de las ondas de radio que llegan a nuestro aparato, disponemos de cierta información adicional, como el hecho de que cada emisora emite en un determinado espectro de frecuencias, y dichos espectros están separados, entonces obviamente el análisis de Fourier permite separar las fuentes. Por eso podemos escuchar la radio y ver la televisión. Por eso se asignan diferentes zonas del espectro electromagnético a los medios de comunicación, a la policía, al ejército, etc.

Ahora bien, si no disponemos de información adicional, o si esta no nos confirma que existe cierta separación entre los espectros de frecuencias utilizados por las distintas fuentes, el análisis de Fourier, por sí mismo, no basta.

Para afrontar esta cuestión, que tiene el nombre técnico de “problema de la separación ciega de fuentes”, es necesario que exista redundancia en la información disponible. Por ejemplo, un mecanismo que permite obtener información redundante consiste en colocar más de un micrófono en la escena. De este modo grabamos la misma señal desde diferentes lugares y, por tanto, sometida a pequeños cambios que nos informarán sobre el origen y forma de las distintas fuentes implicadas. Se trata de algo muy parecido a lo que hacemos constantemente con la vista. Tener dos ojos que están colocados en sitios distintos —pero cuya posición conocemos— nos permite disponer de dos imágenes distintas de la misma escena. Las imágenes son diferentes porque el lugar desde el que mira cada ojo es distinto, y es precisamente esta diferencia la que nos permite “ver” la profundidad. Es lo que conocemos como visión estereográfica. El mismo truco se utiliza para generar imágenes 3D en el cine. Y, en el caso de nuestro sistema auditivo, tener dos oídos nos facilita reconocer la procedencia de los distintos sonidos que percibimos porque, al ser finita —y, de hecho, moderadamente baja— la velocidad de propagación del sonido en el aire, hay un ligero desfase en la recepción del sonido entre un oído y el otro. Este desfase dependerá, por supuesto, del lugar de procedencia del sonido y, por tanto, nos informa sobre este hecho particular.

En el caso del sonido grabado en una habitación en la que se han colocado distintos micrófonos, la información que podemos desgranar va a depender, evidentemente, del número de micrófonos que usemos. A mayor número de micrófonos, tendremos mayor capacidad de discernimiento.

Una vez dispongamos de varios micrófonos registrando el sonido, necesitamos formular matemáticamente tanto la disposición de los micrófonos como los mecanismos físicos que regulan la superposición del sonido que procede de las distintas fuentes sonoras presentes. ¿Hay un único modo válido para describir esta superposición? En otras palabras, antes de abordar el problema de separar las fuentes, debemos concretar cuáles son nuestras hipótesis sobre el proceso de mezcla de las mismas.

En función del contexto en el que se trabaje, existen diversos mecanismos que modelan bien este proceso de mezcla. Evidentemente, los algoritmos de separación de fuentes dependerán del modelo de mezcla que se haya impuesto. En general, si disponemos de  $n$  observaciones realizadas por  $M$  sensores (micrófonos), estas quedan recogidas en una matriz  $X$  del tipo

$$X = \mathbf{col}[x(0), \dots, x(n-1)] = \begin{matrix} & x_1(0) & x_1(1) & \cdots & x_1(n-1) \\ x_2(0) & x_2(1) & \cdots & x_2(n-1) \\ \vdots & \vdots & \ddots & \vdots \\ x_M(0) & x_M(1) & \cdots & x_M(n-1) \end{matrix}$$

Si asumimos que para la generación de los datos registrados han intervenido  $N$  fuentes  $s_i$ ,  $i=1,2,\dots,N$ , estas se modelan con una matriz

$$S = \mathbf{col}[s(0), \dots, s(n-1)] = \begin{matrix} s_1(0) & s_1(1) & \cdots & s_1(n-1) \\ s_2(0) & s_2(1) & \cdots & s_2(n-1) \\ \vdots & \vdots & \ddots & \vdots \\ s_N(0) & s_N(1) & \cdots & s_N(n-1) \end{matrix}$$

y el proceso de mezcla de las fuentes se caracteriza con una matriz

$$A = \begin{matrix} & a_{11} & a_{12} & \cdots & a_{1N} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ & a_{M1} & a_{M2} & \cdots & a_{MN} \end{matrix}$$

y una relación del tipo  $x(k) = A \circ s(k) + e(k)$ ,  $k = 1, \dots, n-1$ , donde  $x(k)$  representa el valor registrado por los micrófonos (los electrodos en el caso de un EEG),  $s(k)$  es el vector de valores de las fuentes, y  $e(k)$  es el ruido introducido en el sistema, en el instante  $k$ . Finalmente,  $\circ$  es la operación de mezcla, cuya naturaleza última dependerá de nuestras hipótesis sobre el modelo. Por ejemplo, si asumimos que se produce una mezcla instantánea, entonces  $\circ$  representa el producto usual de la matriz  $A$  (que recibe el nombre de matriz de mezclas) por el vector columna  $s(k)$ .

En realidad, existen esencialmente dos alternativas claramente diferenciadas en el modelo de mezcla de fuentes: o bien la mezcla se produce de forma instantánea, o bien esta toma en consideración posibles retrasos en la transmisión de las distintas fuentes, retrasos que pueden venir determinados por las diferentes posiciones relativas entre las fuentes y los sensores. Este último caso nos conduce a un modelo basado en convoluciones. Además, las ecuaciones pueden complicarse aún más si tenemos en consideración las posibles reverberaciones (o ecos) que puedan sufrir las señales emitidas por las distintas fuentes. Afortunadamente, si nuestro modelo es convolutivo, la aplicación de la transformada de Fourier traduce las ecuaciones, en el dominio de la frecuencia, a un modelo de mezclas instantáneas. De todas formas, es primordial recalcar que, en el caso más importante que nos ocupa, que es la separación de fuentes a partir de los datos de un electroencefalograma, está bien consolidada la hipótesis de mezclas instantáneas.

Si asumimos el modelo de mezclas instantáneas  $X=AS$ , los vectores fila de la matriz de fuentes  $S$  serán las fuentes. Fijado uno de estos vectores, por ejemplo  $s_k=(s_k(0),s_k(1),\dots,s_k(n-1))$ , el vector correspondiente a la fila  $k$ -ésima de  $S$ , la

columna  $k$ -ésima de la matriz  $A$  nos informará sobre los diferentes pesos con los que esta fuente en concreto contribuye a los datos recogidos por los distintos electrodos del EEG. En efecto, si  $x_i=(x_i(0),\dots,x_i(n-1))$  representa los datos grabados por el  $i$ -ésimo sensor (que están recogidos en la fila  $i$ -ésima de  $X$ ),

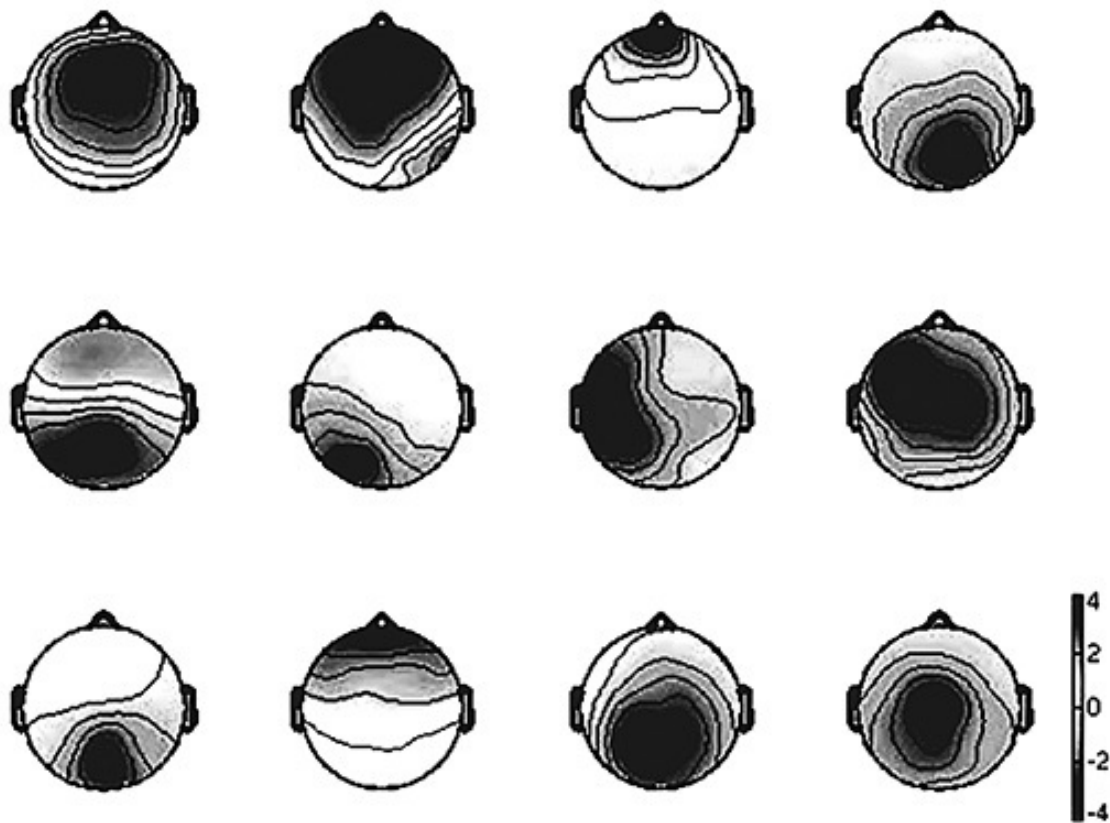
entonces 
$$x_i(t) = \sum_{k=1}^N a_{i,k} s_k(t), \quad \text{para } t = 0, 1, \dots, n-1$$
 y, por tanto,  $a_{i,k}$  representa la contribución de la fuente  $k$ -ésima  $s_k$  al electrodo  $x_i$  (en todo instante de tiempo  $t$ , y para todo sensor  $x_i$ ). A la columna  $k$ -ésima de  $A$  se le asocia un mapa topográfico de la cabeza, que sirve para visualizar las contribuciones al EEG de la componente  $s_k$ . Este mapa se construye tras asignar, para cada valor  $i$ , el peso  $a_{i,k}$  (tomado de la columna  $k$ -ésima de  $A$ ) al electrodo  $x_i$ . La información correspondiente se combina con un método de interpolación (a elegir entre numerosas posibilidades, incluyendo las interpolaciones bilineal y bicúbica, entre otras) que permite asignar un valor a cada uno de los píxeles que aparecen en el topograma que se va a pintar. Sobre dicho topograma, es usual, además, resaltar las curvas de nivel. Finalmente, el topograma se colorea de acuerdo a la regla que determina que a mayor actividad (por ejemplo, cuanto mayor es el valor absoluto número correspondiente) pintamos más oscuro, y a menor actividad, más claro. Si queremos tener en cuenta los signos, una opción habitual es asignar el color blanco al cero, colores amarillos y anaranjados, hasta un rojo o incluso un morado fuerte, a los pesos positivos, y colores verdes y azulados, hasta un azul marino intenso, a los pesos negativos. Para la interpretación de la escala de colores, se coloca una barra lateral que refleja numéricamente las asociaciones que se han realizado. Una simple inspección visual del mapa revela la zona del encéfalo en la que podemos situar la actividad eléctrica de la fuente que estamos investigando.

Esta técnica de construir topogramas se puede usar para visualizar numerosos aspectos del EEG y del ERP, tanto en tiempo como en frecuencia. Incluso existen bases de datos estandarizadas de topogramas contra las que se pueden



contrastar los resultados obtenidos en cada paciente. En estos casos el topograma suele representar la diferencia entre los datos del paciente y un topograma de referencia normalizado, lo cual permite a menudo una interpretación clínica muy eficaz y fácil de explicar.

FIGURA 13  
Algunos ejemplos de mapas topográficos.



Una vez hemos fijado el modelo de mezclas, el siguiente paso es, obviamente, decidir mediante qué mecanismo se van a separar las fuentes. Es decir: ¿cuál es nuestra respuesta al problema de la separación ciega de fuentes? Aquí existen también varias alternativas, cada una de las cuales puede ser muy útil en función del contexto en el que nos encontremos.

De todas las posibles soluciones al problema de separación ciega de fuentes,

en este texto centramos nuestra atención solo en las dos técnicas que han tenido mayor impacto científico en neurofisiología, hasta la fecha. Se trata del análisis de componentes principales y el análisis de componentes independientes (PCA e ICA, respectivamente, en lo que sigue).

## PCA

Supongamos que disponemos de  $M$  sensores y hemos tomado  $n$  muestras y sea  $X \in \mathbf{M}_{M \times n}(\mathbb{R})$  la correspondiente matriz de muestras. Esta matriz siempre admite una descomposición del tipo  $X = U \Sigma V^t$ , donde  $\Sigma$  es una matriz de tamaño  $M \times n$ , cuyas únicas entradas no nulas están en la diagonal y vienen dadas por  $\sigma_{11} := \sigma_1 \geq \sigma_{22} := \sigma_2 \geq \dots \geq \sigma_{rr} := \sigma_r > 0$ , y  $U, V$  son matrices cuadradas que satisfacen, además, las relaciones de ortogonalidad  $UU^t = I_M$ ,  $VV^t = I_n$ . Bajo estas condiciones, decimos que  $X = U \Sigma V^t$  es la descomposición en valores singulares de  $X$  (SVD, en lo sucesivo). Esta descomposición se puede reescribir como

$$X = \sum_{k=1}^r \sigma_k u_k v_k^t, \text{ donde } U = \text{col}[u_1, \dots, u_M] \text{ y } V = \text{col}[v_1, \dots, v_n].$$

En otras palabras, hemos probado que

$$X = \text{col}[\sigma_1 u_1, \dots, \sigma_r u_r] \text{fil}[v_1^t, v_2^t, \dots, v_r^t] = AS,$$

donde  $A = \text{col}[\sigma_1 u_1, \dots, \sigma_r u_r]$  es nuestra matriz de mezclas y  $S = \text{fil}[v_1^t, v_2^t, \dots, v_r^t]$  representa la matriz de fuentes. En este caso, llamamos componentes principales de  $X$  a los vectores  $v_1^t, v_2^t, \dots, v_r^t$  y a la descomposición así obtenida la llamamos descomposición en componentes principales o PCA de  $X$ .

Existen varios algoritmos que calculan la descomposición en componentes principales con exquisitas velocidad y precisión. De hecho la SVD es uno de los objetos matemáticos “estrella” en el mundo del álgebra lineal numérica, y ha sido (y es) estudiada en profundidad a lo largo de las últimas décadas. En

particular, la descomposición en componentes principales se ha considerado especialmente importante en numerosas aplicaciones porque tiene algunas propiedades matemáticas muy interesantes. Por ejemplo, se sabe que, para todo  $k$

$\leq r$ , la matriz  $B_k = \sum_{i=1}^k \sigma_i u_i v_i^t$  tiene rango  $\leq k$  y, entre todas las matrices  $B$  de rango  $\leq k$ , minimiza  $\|X - B\|_F$ , donde  $\|A\|_F = \sqrt{\sum_{i,j} A_{ij}^2}$  denota la norma Euclídea de la matriz  $A$ , verificándose, además, la identidad  $\|X - B_k\|_F = \sigma_{k+1}$ . Aunque esta y otras propiedades relacionadas son razón más que suficiente para motivar el estudio matemático de la SVD y del PCA, lo cierto es que las fuentes calculadas de este modo, a partir de la matriz de muestras  $X$  obtenida de un EEG, no admiten una interpretación fisiológica clara. En particular, no es razonable asumir que los vectores columna de la matriz de mezclas, que definen los correspondientes mapas topográficos y dan, por tanto, cuenta del origen físico de las distintas fuentes, deban formar necesariamente conjuntos de vectores ortogonales. Aun así, el PCA se utiliza para el preprocesado de los datos, entre otras cosas.

## ICA

A mediados de la década de 1990 los neurobiólogos computacionales A. J. Bell y T. J. Sejnowski introdujeron una nueva técnica de separación de fuentes que, desde un punto de vista fisiológico, está mucho más justificada que el PCA. Se trata del análisis de componentes independientes o ICA. La idea es buscar fuentes que, interpretadas como variables aleatorias, posean la propiedad de explicar los datos y ser estadísticamente independientes.

En efecto, parece muy apropiado asumir que la información que proviene de las distintas fuentes que generan los potenciales eléctricos recogidos en el EEG actúan de forma independiente. Diferentes redes neuronales se activan en diferentes partes del encéfalo —que podemos considerar fijas a lo largo de la grabación del EEG— para dar respuesta a los estímulos que reciben, o por

iniciativa propia, y estas actúan independientemente unas de otras. Estas son las fuentes que nos gustaría ser capaces de localizar y describir. Además, la información que provenga de artefactos debe considerarse también estadísticamente independiente de la información que responde a la actividad del encéfalo que, ciertamente, es lo que queremos medir.

En lo que queda de esta sección nos ocupamos de proporcionar la descripción formal del ICA en términos matemáticos. En la próxima sección explicaremos cómo se aplica esta técnica para eliminar artefactos del EEG.

El punto de partida, como en muchos otros problemas en los que la teoría de probabilidades interviene de forma significativa, es el llamado Teorema Central del Límite. Dicho resultado afirma, en pocas e imprecisas palabras, que en determinadas circunstancias, cuando se produce la suma de un elevado conjunto de procesos aleatorios que son, además, estadísticamente independientes, el resultado se acerca cada vez más (cuando crece el número de sumandos) a un proceso aleatorio gaussiano. Los matemáticos dicen, para resumir la situación, que la suma de variables aleatorias independientes (e idénticamente distribuidas) tiende a comportarse como una variable muy concreta, que recibe el nombre de variable normal o gaussiana. Estas variables se rigen por la ley

$$P[a \leq X \leq b] = \frac{1}{\sqrt{2\pi\sigma^2}} \int_a^b e^{-\frac{(t-\mu)^2}{2\sigma^2}} dt,$$
 donde  $\mu$  y  $\sigma^2$  representan la media y la varianza de la variable aleatoria  $X$ . El Teorema Central del Límite afirma que si tomamos una sucesión de variables aleatorias independientes e idénticamente distribuidas

$X_k$ ,  $k=1,2,\dots$ , y consideramos la sucesión de sumas  $S_N = \sum_{k=1}^N X_k$ , entonces, si  $\mu$  y  $\sigma^2$  denotan la media y la varianza de las variables  $X_k$ , la sucesión de variables

aleatorias  $Z_N = \frac{S_N - N\mu}{\sigma\sqrt{N}}$  converge a una variable aleatoria normal con media 0 y varianza 1, cuando  $N$  tiende a infinito. Concretamente, son sus funciones densidad de probabilidad las que convergen a la función densidad de

probabilidad de la normal de media 0 y varianza 1, que está dada por

$$p_{N(0,1)}(t) = \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}}.$$

Es probable que algún lector se haya perdido justo en este momento. El resultado parece demasiado técnico, ¿no? Requiere usar numerosos conceptos que aún no hemos introducido, como los de variable aleatoria, distribución, convergencia, media, varianza, función densidad de probabilidad, variable aleatoria normal, independencia estadística, etc. ¿Qué es todo esto? Antes de proseguir, vamos a describir brevemente cada uno de estos términos.

Cuando realizamos una experiencia cuyo resultado podemos medir pero no predecir de forma completamente precisa, estamos en un contexto que debe ser analizado desde el punto de vista de la probabilidad. Es decir, en este tipo de situaciones interviene el azar. No podemos adelantar cuál será el resultado del experimento. Lo que sí podemos hacer, en muchos casos, es algún tipo de afirmación probabilística. Por ejemplo, podemos decir cosas como: “Con probabilidad superior a 1/2, el valor obtenido será positivo”. De hecho, esto resulta muy útil en numerosas circunstancias, y la teoría de probabilidad es la herramienta matemática encargada de manejar este tipo de situaciones. Es decir, aunque parezca contradictorio, resulta que el azar tiene leyes y estas son, además, de enorme utilidad práctica.

Los voltajes que medimos en cada electrodo de un EEG son, en términos generales, funciones que oscilan de manera azarosa. Esto no significa que no respondan a razones precisas, sino que son el resultado de la suma de muchísimos procesos diferentes, actuando simultáneamente. En estas circunstancias es imposible predecir de forma determinista el valor que se obtendrá en un electrodo en un instante determinado y, por tanto, un enfoque basado en el uso de probabilidades es necesario. Del mismo modo, resulta imposible determinar a priori el valor de otros muchos parámetros biológicos, como el peso o la altura de una persona en función de su edad, sexo y régimen alimenticio. Sin embargo, sí podemos determinar ciertos intervalos en los que

dichos valores se van a situar en la mayoría de los sujetos, porque estas variables tienen un comportamiento probabilístico conocido.

Podemos, por tanto, pensar en los potenciales eléctricos medidos por el EEG como resultado de tomar muestras de los valores que se obtienen tras realizar un experimento aleatorio repetidas veces, siempre el mismo experimento. Cuando hacemos esto, decimos que estamos considerando una variable aleatoria (v.a.). Cada electrodo define una v.a. y los valores obtenidos por un electrodo en un EEG representan un conjunto de muestras de dicha v.a.

Si denotamos por  $X$  una v.a. concreta, ¿cómo podemos estudiar su comportamiento? Le asociamos una función, llamada distribución de probabilidad de la variable, que la define en términos probabilísticos. Se trata de la función  $F_X(t) = P[X \leq t]$ , donde  $P[X \leq t]$  denota la probabilidad de que, al tomar una muestra de la variable aleatoria  $X$ , el valor obtenido sea inferior al número real  $t$ . Estas funciones nos proporcionan toda la información necesaria para tratar con las variables aleatorias que las definen. Con su ayuda se pueden hacer todo tipo de inferencias estadísticas y son, de hecho, una poderosa herramienta de trabajo.

En la práctica existen esencialmente dos tipos de variables aleatorias, que están perfectamente diferenciados: las variables continuas y las discretas. Diremos que la variable aleatoria  $X$  es discreta si existe una sucesión

(posiblemente finita) de números  $\{y_k\}_{k=1}^{\infty}$ , que llamaremos rango de la variable, tal que la variable aleatoria solo puede tomar valores que pertenezcan a dicha sucesión. Cuando esto es así, la distribución de probabilidad de la variable

aleatoria  $X$  queda determinada por la expresión:

$$P[X \leq t] = \sum_{y_k \leq t} P[X = y_k] \quad \text{y, por}$$

tanto, los números  $p_k := p_X(y_k) = P[X = y_k]$ ,  $k = 1, 2, \dots$ , determinan con toda precisión su comportamiento. Si la v.a. es continua, entonces la función de distribución  $F_X(t)$  viene determinada típicamente como una integral

$F_X(t) = P[X \leq t] = \int_{-\infty}^t p_X(s) ds$  para cierta función positiva  $p_X(t)$  que recibe el nombre de función densidad de probabilidad de  $X$  (fdp, en lo sucesivo). Por supuesto, también existen variables mixtas y a veces es necesario considerar la integral anterior en el sentido de la teoría de la medida. Finalmente, si consideramos un vector  $(X_1, \dots, X_m)$  formado por v.a., este admite también una descripción probabilística en términos similares al caso univariado. En particular, si se trata de v.a. continuas, su fdp  $p_{(X_1, \dots, X_m)}$  queda determinada por la igualdad

$$P[X_1 \leq t_1, \dots, X_m \leq t_m] = \int_{-\infty}^{t_1} \dots \int_{-\infty}^{t_m} p_{(X_1, \dots, X_m)}(s_1, \dots, s_m) ds_1 \dots ds_m.$$

Aunque hay infinidad de distribuciones de probabilidad, solo un reducido grupo de ellas basta para estudiar la mayor parte de fenómenos aleatorios que encontramos en la práctica. De hecho, cada rama del saber tiene asociadas unas pocas distribuciones de probabilidad cuyo conocimiento basta para entender buena parte de los fenómenos que estudia. Las más habituales son, entre las discretas, las distribuciones binomial (definida por

$$P[Bin(N, p) = k] = \binom{N}{k} p^k (1-p)^{N-k} \quad \text{y de Poisson (definida por } P[Pois(\lambda) = k] = \frac{e^{-\lambda} \lambda^k}{k!} \text{)}.$$

En el caso de las v.a. continuas, las más importantes son la normal (cuya fdp es

$$p_{N(\mu, \sigma^2)}(t) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(t-\mu)^2}{2\sigma^2}} \text{), la uniforme (cuya fdp está dada por } p_{U(a,b)}(t) = \frac{1}{b-a} \text{ para } t \in [a,b] \text{ y } p_{U(a,b)}(t) = 0 \text{ para } t \notin [a,b] \text{) y la exponencial (cuya fdp es } p_{Exp(\lambda)}(t) = \lambda e^{-\lambda t} \text{ para } t \geq 0 \text{ y } p_{Exp(\lambda)}(t) = 0 \text{ para } t < 0 \text{)}$$

Dos v.a.  $X, Y$  se dicen independientes si para todo par de valores  $t, s \in \mathbb{R}$ , se tiene que  $P[X \leq t \wedge Y \leq s] = P[X \leq t]P[Y \leq s]$ , lo cual significa, en términos sencillos, que el valor que toma cualquiera de las v.a. no influye en el correspondiente valor de la otra. Recuérdese que dos sucesos probabilísticos  $A$ ,

$B$  serán independientes si la probabilidad de que ambos se den es precisamente el producto de las probabilidades de cada uno de ellos:  $P(A \cap B) = P(A)P(B)$ . Por ejemplo, si lanzamos dos veces un dado (no trucado), la probabilidad de que en el primer lanzamiento obtengamos 1 y en el segundo obtengamos 6 es  $1/36 = (1/6)^2$  porque el resultado del primer lanzamiento no condiciona para nada el resultado del segundo lanzamiento (el dado no tiene memoria).

En el caso de que las v.a.  $X, Y$  sean continuas, su independencia se traduce en que sus fdp satisfacen la relación  $p_{(X,Y)}(t,s) = p_X(t)p_Y(s)$ . En el caso de que las v.a. sean discretas con rangos  $\{x_k\}$  e  $\{y_s\}$ , respectivamente, la independencia se traduce en la relación  $P[X = x_k \wedge Y = y_s] = P[X = x_k]P[Y = y_s]$ . La definición de v.a. independientes se extiende de manera obvia a cualquier número finito de variables.

Asociadas a una v.a. existen ciertas magnitudes que tienen interés porque resumen, en un solo número, algunos aspectos importantes de la variable. Las más importantes son la esperanza matemática y la varianza. Además, son también de interés los llamados momentos de orden superior y, cuando hay varias variables en juego, las correlaciones y covarianzas. El momento de orden  $n$  es el valor

$$E(X^n) = \begin{cases} \int_{-\infty}^{\infty} t^n p_X(t) dt & \text{si } X \text{ continua} \\ \sum_{k=1}^{\infty} x_k^n P[X = x_k] & \text{si } X \text{ discreta con rango } \{x_k\}_{k=1}^{\infty} \end{cases}$$

Cuando  $n=1$ ,  $E(X)$  se llama esperanza, valor esperado, o valor medio de  $X$ . La varianza de  $X$  es  $V(X) = E((X - E(X))^2) = E(X^2) - E(X)^2$ . La covarianza de dos v.a.  $X, Y$  es  $cov(X, Y) = E(XY) - E(X)E(Y)$ , y representa una medida de “correlación” entre ambas variables. Si las v.a.  $X_1, \dots, X_m$  son independientes, entonces

$$cov(X_i, X_j) = E(X_i X_j) - E(X_i)E(X_j) = 0 \text{ siempre que } i \neq j,$$



es decir, las v.a. no están correlacionadas. El enunciado inverso no es cierto: la independencia es mucho más exigente que la ausencia de correlación. Otra medida que es especialmente importante para nuestros intereses, es la curtosis,

que se calcula mediante la fórmula 
$$K(X) = \frac{E((X - E(X))^4)}{V(X)^2} - 3.$$
 Si la v.a.  $X$  es normal entonces su curtosis vale 0. Por tanto, una forma de medir si una v.a. concreta está cerca o lejos de ser gaussiana es calcular su curtosis. Si es diferente de 0, la v.a. no será gaussiana y lo será tanto menos cuanto mayor sea su módulo.

Ahora que hemos introducido todos estos conceptos, podemos volver al problema de la separación ciega de fuentes. Cada electrodo toma muestras del potencial eléctrico con cierta frecuencia (todos con la misma frecuencia de muestreo, evidentemente, porque nos interesa que estén coordinados en el tiempo). Esto significa que, si se han tomado  $n$  muestras, y suponemos que el modelo de mezcla es lineal, instantáneo y sin ruido, tendremos una relación del tipo  $x(k) = A \cdot s(k)$   $k = 1, \dots, n-1$ , donde  $x(k)$  es el vector columna que representa el valor registrado por los electrodos en el instante de tiempo  $k$  (es decir, en la  $k$ -ésima muestra) y  $s(k)$  contiene la información de las fuentes originales, en el mismo instante de tiempo  $k$ . Por su parte, la matriz de mezclas  $A$  nos da cuenta de cómo se han superpuesto estas fuentes para producir la información que hemos grabado. Para el problema que estamos considerando, vamos a imponer que el número de fuentes  $s_i$  nunca supere al número de electrodos  $y$ , de hecho, es usual asumir que ambos números coinciden. Posteriormente, durante la aplicación del algoritmo, puede suceder que este nos indique que el número de fuentes es en realidad menor que el número  $N$  de sensores. Esencialmente, las afirmaciones anteriores responden al principio físico que garantiza que si solo tenemos  $N$  micrófonos, no podremos separar más de  $N$  sonidos independientes. Con un solo micrófono no podemos separar sonidos, a no ser que dispongamos de información extra, como la separación en frecuencias en anchos de banda disjuntos, en cuyo caso el análisis de Fourier

permite hacer los cálculos pertinentes (pero este no es el caso del problema que nos ocupa). Con dos micrófonos podremos distinguir hasta dos fuentes distintas, con tres, tres, etc.

Por tanto, en nuestro caso estamos imponiendo que la matriz  $A$  sea cuadrada. Si conocemos  $A$  y, además, sabemos que es invertible, con inversa  $W=A^{-1}$ , entonces las fuentes originales se calculan simplemente con la fórmula  $s(k) = W \cdot x(k)$ ,  $k = 1, \dots, n-1$ . En otras palabras, para resolver el problema de separación de fuentes lo único que, de verdad, necesitamos, es calcular la matriz  $W$ . ¿Cómo se hace esto? Todo depende del modelo de separación de fuentes al que queramos acogernos. No existe una única forma de descomponer los datos que hemos grabado como resultado de superponer distintas “fuentes”. De hecho, existen infinitas formas de hacer esto, ¡pero solo una puede corresponderse con la “realidad”! En efecto: de lo que se trata aquí es conocer qué hipótesis sobre las fuentes (además de las que ya hemos realizado sobre el proceso de mezcla de las mismas, al cual hemos impuesto la linealidad y la mezcla espontánea) son razonables para que nuestro modelo matemático refleje, lo mejor que se pueda, el fenómeno físico que pretende describir.

Pensando en esto, es muy natural suponer, como ya hemos advertido anteriormente, que los diferentes mecanismos que actúan en nuestro encéfalo para producir los potenciales eléctricos que estamos midiendo (sean los que sean), se comportan como procesos independientes. Esta hipótesis es el fundamento del modelo ICA y, afortunadamente, se puede aprovechar matemáticamente porque el Teorema Central del Límite nos garantiza que si diversas fuentes independientes se superponen para producir la actividad en un electrodo, la v.a. que representa dicho electrodo debe estar más cerca de la normalidad que cualquiera de las fuentes originales. Por tanto, si conocemos una medida de “gaussianidad” (y la curtosis es una de dichas medidas), podemos utilizar esta información para buscar las fuentes que originan nuestros datos.

Para aplicar este método es necesario asumir que las fuentes originales no se

corresponden con procesos gaussianos, pero resulta gratificante comprobar que precisamente esta hipótesis es también natural en el contexto de la electroencefalografía. Así las cosas, lo que queda es describir cómo atacar el problema desde un punto de vista estrictamente matemático.

Por construcción, si  $X = \text{fil}[x_1(t), x_2(t), \dots, x_N(t)]$  denota la matriz formada por los valores registrados por los  $N$  electrodos, y  $S$  es la matriz formada por las fuentes, entonces  $S = \text{fil}[s_1(t), s_2(t), \dots, s_N(t)] = W \text{fil}[x_1(t), x_2(t), \dots, x_N(t)]$  (la variable  $t$  que aparece en  $x_i(t)$  y  $s_i(t)$  representa el tiempo. Así,  $x_i(t) = [x_i(0), x_i(1), \dots, x_i(n-1)]$ ). Se sigue que, si  $W = \text{fil}[w_1^T, \dots, w_N^T]$ , entonces

$$s_1(t) = w_1^T \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_N(t) \end{pmatrix}.$$

Por tanto, para calcular la primera componente  $S_1$ , debemos buscar un vector de pesos  $w_1^T = [w_{11}, \dots, w_{1N}]$  tal que el nuevo vector  $s_1 = [s_1(0), s_1(1), \dots, s_1(n-1)]$ , que viene dado por la expresión

$$s_1(t) = \sum_{k=1}^N w_{1k} x_k(t), \quad 0 \leq t \leq n-1$$

, alcance el máximo valor absoluto de su curtosis. Es decir, nos enfrentamos a un problema de optimización para una función que

depende de  $N$  variables,  $F(w_{11}, \dots, w_{1N}) = \left| K \sum_{k=1}^N w_{1k} x_k \right|$ .

Además, como el valor de la

curtosis de la mezcla  $\sum_{k=1}^N w_{1k} x_k$  no depende del módulo del vector  $w_1^T$  (lo cual se comprueba mediante un sencillo cálculo), resulta que podemos imponer que

dicho vector pertenezca a la esfera de radio 1 de  $\mathbb{R}^N$ . La compacidad de la esfera garantiza que existe el máximo buscado. De hecho, este resultado, garantizando que si  $f : S \rightarrow \mathbb{R}$  es continua y  $S$  es cerrado y acotado en  $\mathbb{R}^N$  (para algún  $N \in \mathbb{N}$ ) entonces  $f$  alcanza su máximo absoluto en  $S$  (y también su mínimo absoluto), fue demostrado ya en el siglo XIX por el matemático alemán K. Weierstrass, y es el caso más sencillo donde podemos resolver problemas de optimización. La técnica del análisis matemático que se utiliza con más frecuencia para buscar uno de estos puntos extremos (y, por tanto, una de las componentes independientes), es el llamado método del descenso del gradiente, y se ha comprobado ampliamente su eficacia. Hay otras técnicas, basadas en la búsqueda de puntos fijos, para lo cual existe una amplia batería de resultados matemáticos disponibles. En particular, uno de los algoritmos que se utilizan con más frecuencia, el denominado FastICA (debido a la velocidad con la que opera, así como a su estabilidad), se basa precisamente en técnicas de punto fijo. Típicamente, existen  $2N$  soluciones para este problema, pero esencialmente son  $N$  y cambios de signo. Cuando hallamos cualquiera de ellas, se pueden obtener las otras imponiendo que, además de resolver el problema de optimización, las nuevas fuentes no estén correlacionadas con las soluciones halladas anteriormente. De esta forma se calculan las  $N$  componentes independientes asociadas al EEG.

Es importante observar, sin embargo, que la curtosis, calculada a partir de una muestra de una v.a., no es una medida robusta de no-gaussianidad. Su valor puede depender de unas pocas medidas erróneas o irrelevantes y, por tanto, aunque esta medida nos ha servido para introducir de forma muy pedagógica el método ICA, en realidad se hace imprescindible aplicar este tipo de algoritmo basándonos en otras medidas de gaussianidad cuya robustez no esté en duda. La más conocida, se basa en las propiedades de la entropía y la información.

Para explicar qué entendemos por información, planteamos la siguiente experiencia: lanzamos al aire una moneda (no trucada). ¿Qué información nos da

el resultado de dicha experiencia, si sale, por ejemplo, cruz? Nos proporciona exactamente un bit de información. De hecho, por definición, un bit es la cantidad de información obtenida al especificar una de dos posibles alternativas que son igualmente probables (por ejemplo, cara o cruz, en el lanzamiento de una moneda no trucada). Si al realizar un experimento aleatorio se produce el suceso  $A$ , cuya probabilidad es  $P(A)$ , diremos que dicho suceso ha aportado

$I(A) = \log_2 \frac{1}{P(A)}$  bits de información sobre el experimento. Nótese que si  $P(A)=1/2$ , entonces  $1/P(A)=2$  y  $\log_2(1/P(A))=\log_2 2=1$ , que es lo que queríamos. Por otra parte, dado un experimento aleatorio  $\Upsilon$ , cabe preguntarse cuál es la información media que proporciona cada realización del mismo. Esta magnitud es especialmente interesante para las aplicaciones y recibe el nombre de entropía.

Si el experimento se traduce en evaluar una v.a.  $X$ , entonces estaremos considerando la entropía de la v.a.,

$$H(X) = \begin{cases} \bullet \sum_{k=0}^{\infty} P[X = x_k] \log_2 \frac{1}{P[X = x_k]} & \text{si } X \text{ v.a. discreta} \\ \equiv \int_{-\infty}^{\infty} p_X(t) \log_2 \frac{1}{p_X(t)} dt & \text{si } X \text{ v.a. continua} \end{cases}$$

En el caso de las v.a. discretas, la entropía es siempre una cantidad positiva, puesto que las probabilidades  $P[X=x_k]$  son siempre a lo sumo 1 y, por tanto,

$\log_2 \frac{1}{P[X = x_k]} = -\log_2 P[X = x_k] \geq 0$ . Sin embargo, cuando la v.a. es continua, bien puede suceder que su función densidad de probabilidad  $p_X(t)$  tome valores superiores a 1 y, en ocasiones, la entropía resultante puede alcanzar un valor negativo. En todo caso, la entropía es una magnitud que refleja el nivel de “predecibilidad” asociado a la variable aleatoria (o lo que es lo mismo, al experimento aleatorio bajo consideración).

A menudo se considera la “entropía negativa” o “sintropía”, que es la cantidad  $J(X) = H(X_{gauss}) - H(X)$ , donde  $X_{gauss}$  es una v.a. gaussiana que tiene las mismas media y varianza que la v.a.  $X$ . La sintropía toma siempre valores positivos y se hace 0 cuando la v.a.  $X$  es gaussiana. Así pues, será mayor cuanto más alejada esté la v.a. de ser gaussiana y, por tanto, un buen criterio para la búsqueda de componentes independientes es la maximización de la sintropía de  $S=WX$ .

Otra magnitud que tiene especial interés es la información mutua. Si  $X = (X_1, \dots, X_N)$  con  $X_i$  v.a. escalares, entonces la información mutua de las v.a.  $X_1, \dots, X_N$

viene dada por 
$$I(X_1, \dots, X_N) = -H(X) + \sum_{k=1}^N H(X_k),$$
 donde  $H(X)$  se calcula en términos de la distribución de probabilidad conjunta de las v.a.  $X_1, \dots, X_N$ . Se puede demostrar que  $I(X_1, \dots, X_N)$  es una cantidad no negativa, y que se anula solo en el caso de que las v.a. sean independientes. Por tanto, podemos construir un algoritmo ICA mediante la búsqueda de una matriz  $W$  tal que  $S=WX$  minimice la información mutua. Obsérvese que en este caso no estamos comprobando la no-gaussianidad de las fuentes sino que aplicamos directamente un criterio de independencia estadística. Este criterio nos muestra que, como es natural, existen diversas formas de detectar que varias v.a. son independientes, y cualquiera de ellas es susceptible, en principio, de ser usada para la construcción de un algoritmo ICA.

El algoritmo que acabamos de describir tiene diversas variantes. Fue introducido en 1995 por A. J. Bell y T. J. Sejnowski bajo el nombre de InfoMax, el cual se justifica porque logran la minimización de la información mutua

mediante la maximización de la entropía 
$$H(X) = -I(X_1, \dots, X_N) + \sum_{k=1}^N H(X_k).$$

Se aplica bajo la hipótesis de que las fuentes son todas supergaussianas o todas subgaussianas, pero funciona mejor bajo la hipótesis de supergaussianidad. Posteriormente, fue mejorado para cubrir todos los casos.

Es de vital importancia observar que algunas de las magnitudes con las que estamos trabajando, como la entropía o la información mutua, no se conocen directamente a partir de la observación de las muestras, sino que dicha información se utiliza para, mediante el uso de estadísticos apropiados, obtener estimaciones razonables de su verdadero valor. Es decir, en este proceso entra en juego de forma inevitable y como parte fundamental en todos los algoritmos que podamos utilizar, la teoría de la estimación estadística. Un aspecto especialmente interesante porque aún no se ha logrado un argumento teórico definitivo que cierre la cuestión, es averiguar cuántas muestras son necesarias para garantizar que la descomposición en componentes independientes se logra de forma exitosa (es decir, realmente separa los datos como superposición de fuentes estadísticamente independientes). La práctica nos dice que, para un número bajo de electrodos, del orden de 20, basta tomar 10.000 muestras por electrodo (y, por supuesto, al recoger un EEG se toman muchísimas más). Pero es un problema abierto encontrar las condiciones precisas para que esto se garantice, en general. Otro aspecto interesante que hay que resaltar es que la introducción de filtros apropiados, como parte del preprocesado anterior a la aplicación del algoritmo ICA, facilita generalmente una buena separación de fuentes.

Por todo lo dicho anteriormente, queda claro que no hay un único algoritmo que descomponga el conjunto de señales en términos de sus componentes independientes, sino que, por el contrario, existe toda una batería de algoritmos que persiguen dicho objetivo. De hecho, existen distintos paquetes de *software* que utilizan algoritmos distintos y, por supuesto, los resultados pueden variar ligeramente de unos a otros. Lo importante es que el modelo teórico nos garantiza que cualquiera de estos algoritmos proporciona una descomposición fiable, la cual puede usarse con toda garantía en las distintas aplicaciones existentes. Esto se debe a que, en cierto modo, el problema al que nos enfrentamos es estable. Es decir, los mismos datos, sometidos a algoritmos ICA distintos, producen resultados muy similares, y datos parecidos sometidos a cualquier algoritmo ICA producen resultados también parecidos. Si hacemos,

por ejemplo, pruebas con sonido —separando la señal grabada por varios micrófonos cuando varias fuentes sonoras actúan simultáneamente— podremos hacernos una idea bastante precisa de lo bien que funcionan todos estos algoritmos. ¡Resuelven, con técnicas estadísticas, un problema al que el todopoderoso análisis de Fourier no puede enfrentarse! ¿Cuáles son, pues, los criterios que debemos tener en consideración a la hora de seleccionar un algoritmo ICA concreto? Lo ideal es buscar, entre todos los existentes, aquel que nos proporcione mayor estabilidad y mayor velocidad de cómputo. Téngase en cuenta que los EEG son matrices de dimensiones gigantescas y, por tanto, debemos buscar también la eficiencia en el cálculo. Nosotros utilizamos, para las aplicaciones en diagnóstico de trastornos mentales, el algoritmo InfoMax, tal como viene implementado en el paquete informático WinEEG, desarrollado por Y. Kropotov en Rusia y que se basa en los artículos originales de Bell y Sejnowski de 1995 y 1999. Otro algoritmo frecuentemente utilizado en el estudio de EEG, que ya hemos mencionado anteriormente, es FastICA, el cual fue desarrollado por A. Hyvärinen en Finlandia.

Tras leer los párrafos anteriores, no sería extraño que algunos lectores empezaran a dudar de la eficacia del método que estamos describiendo, ya que su aplicación requiere, en principio, asumir numerosas hipótesis que podrían no ser completamente satisfechas en los ejemplos a los que posteriormente debemos enfrentarnos. Piénsese que, después de todo, los electroencefalogramas parecen oscilar de forma tremendamente azarosa, y aquí estamos afirmando que, aunque la filosofía de trabajo es la misma, en general, existen diversas formas de implementar, dando lugar a algoritmos distintos, con resultados que podrían, por tanto, variar. Además, en la ejecución final de los mismos, cuya base es una mezcla de estadística y análisis, se requiere aproximar algunos datos esenciales, como la entropía diferencial o la curtosis, mediante el uso de estadísticos. ¿No será esto excesivo? ¿No estaremos corriendo demasiados riesgos cuando, además, deseamos aplicar estos cálculos a datos verdaderamente sensibles y en pacientes humanos? Procede, por tanto, hacer una defensa que elimine las dudas.



Es fundamental que se registren los EEG con exquisito cuidado (y esto es algo que no debemos cansarnos de repetir), pues los fallos en este punto del proceso invalidarán su posible uso clínico. Además, cuando generamos la base de datos estándar contra la cual vamos a realizar nuestros contrastes y, por tanto, gracias a la cual podremos hacer diagnósticos y protocolos de actuación, los EEG deben ser recogidos y analizados siempre del mismo modo, con el mismo *software* y *hardware* y los mismos algoritmos de cálculo. Es esta precisamente la garantía de su utilidad. Aun así, podríamos persistir en nuestras dudas debido a que se requiere el uso de estadísticos que aproximen las medidas teóricas en las que se basan los algoritmos y, además, sería necesario conocer algunos resultados sobre la estabilidad de dichos procesos.

Cuando tenemos la ocasión de explicar el algoritmo ICA de forma presencial, ante un público atento, resulta muy natural introducir un ejercicio práctico en el que lo aplicamos a una serie de sonidos que grabamos y mezclamos directamente en el aula, con participación de los presentes. Tras mezclar los sonidos y reproducir dichas mezclas —que son bastante ininteligibles, por lo general—, realizamos (en breves segundos) el proceso de separación de fuentes, con FastICA. Luego escuchamos el resultado. Todo el mundo se convence de que la separación se ha logrado con enorme calidad y rapidez. Generalmente hacemos esta práctica con la intervención de tres o cuatro personas que se prestan voluntariamente. Cada uno lee un poema, que se graba a 8.000 Hz y cuya lectura tiene una duración máxima de entre diez y quince segundos. A continuación mezclamos las fuentes multiplicando por una matriz de mezclas invertible, escuchamos los micrófonos, y a ellos se les aplica el algoritmo ICA. Finalmente, reproducimos el sonido correspondiente a cada una de las fuentes halladas por el algoritmo. ¡El resultado es siempre excelente! ¡El público se emociona!

Mostramos ahora de forma visual cuáles son los efectos del algoritmo ICA aplicado a imágenes en escala de gris. Este tipo de objetos quedan definidos a partir de una matriz de valores. Cada entrada representa el nivel de gris que se va

a utilizar para pintar el pixel de la imagen que ocupa la posición de la entrada de la matriz que se ha seleccionado. Por tanto, para poder aplicar el algoritmo necesitamos en primer lugar transformar la matriz que define una imagen en escala de grises, en un vector. Esto se logra leyéndola por filas y apuntando de forma ordenada los registros correspondientes en un único vector. Si tenemos varias imágenes que se supone son mezcla de otras, aplicamos ICA a los vectores obtenidos a partir de dichas matrices del modo que acabamos de indicar. Luego cada uno de los vectores hallados por ICA se transforma en una matriz mediante el mecanismo inverso al que se utilizó para su formación. Es decir, rompiendo en trozos el vector (cada uno de ellos del tamaño de las filas de la imagen original) y colocando los trozos uno tras otro para formar la correspondiente matriz imagen. Finalmente, es necesario normalizar la matriz resultante, tomando su módulo y multiplicando por un escalar apropiado, para que el ordenador pueda mostrarla. Los resultados son también bastante buenos, como se puede observar en el ejemplo que incluimos aquí (véanse las figuras siguientes). Es evidente que las componentes independientes halladas de esta forma reflejan bastante bien las auténticas fuentes originales a partir de las cuales hemos generado las mezclas. Captan, sin lugar a dudas, lo esencial, y eso debe animarnos.

**FIGURA 14**  
Imágenes mezcladas.



FIGURA 15  
Fuentes calculadas con FastICA.



## Eliminación de artefactos en electroencefalografía

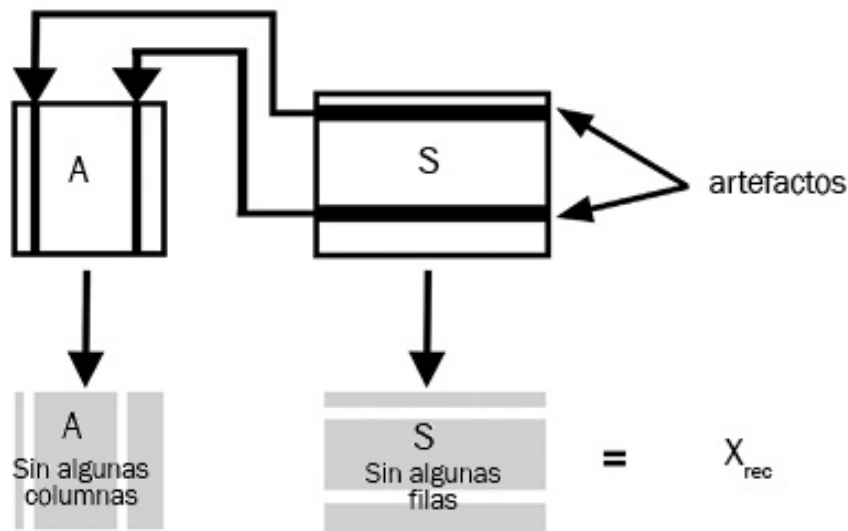
Cuando grabamos el EEG, es inevitable que aparezcan artefactos. Esto es sencillo de comprobar en la práctica si disponemos de un sensor cualquiera, capaz de registrar la señal EEG desde el cuero cabelludo, y de algún tipo de

*software*, conectado al sensor (por ejemplo, existen aplicaciones sencillas para teléfonos móviles que hacen esta operación), que muestre en tiempo real la señal recogida por el sensor. Basta parpadear para observar que el EEG registrado se modifica sustancialmente en el momento del parpadeo, y que cuando este cesa, vuelve a mostrar una señal similar a la que se podía observar con anterioridad. Esto significa, obviamente, que actividades sencillas, muchas veces inconscientes, como el parpadeo, pueden generar potenciales eléctricos que luego son recogidos por el EEG. Evidentemente, nos gustaría eliminar de la señal EEG que vamos a estudiar la influencia del parpadeo, así como otras posibles perturbaciones provocadas, por ejemplo, por la actividad cardíaca o el movimiento de los músculos faciales. A este proceso se le denomina desartefactado del EEG y es, claro está, algo que debemos hacer siempre que grabemos un EEG, antes de someterlo a estudio. El parpadeo, así como los movimientos oculares, durante mucho tiempo se han eliminado mediante la colocación de un sensor cerca de cada ojo, que registra este tipo de actividad, y el uso posterior de dicha señal para eliminarla del EEG. Esto, sin embargo, resulta desagradable para los pacientes, y en la actualidad es una tarea que se realiza aplicando ICA a la señal registrada por el EEG. La razón principal por la que sabemos que podemos utilizar ICA para el desartefactado es que las fuentes que originan los artefactos reflejan necesariamente procesos estocásticos independientes de la actividad eléctrica que deseamos medir. Es decir, las respuestas a los estímulos o incluso la actividad eléctrica espontánea es una actividad necesariamente independiente de la actividad asociada a los artefactos, como el parpadeo. Por tanto, el algoritmo ICA tiene la capacidad de separarlos. Para decidir si una componente determinada, hallada mediante ICA, se corresponde con uno o varios artefactos y, por tanto, debe ser eliminada del EEG, miramos su mapa topográfico asociado. Dicho mapa nos informa sobre la localización física de la fuente y nos permite, de paso, decidir, en base a criterios clínicos, si se trata o no de un artefacto.

Supongamos que hemos aplicado ICA para obtener la descomposición  $X=AS$ ,

donde  $A$  es la matriz de mezclas y  $S$  es la matriz de fuentes. Las filas de  $S$  representan las distintas fuentes halladas. Para cada componente independiente  $s_i$  que hemos identificado como un artefacto, se procede a su eliminación del EEG.

FIGURA 16  
Desartefactado del EEG.



¿Cómo se hace esto? Basta borrar la fila  $i$ -ésima de  $S$  y la columna  $i$ -ésima de  $A$ . Si  $S^*$  y  $A^*$  denotan las matrices que se obtienen tras realizar las correspondientes supresiones, entonces  $X_{rec} = A^* S^*$  representará el EEG reconstruido, sin artefactos. Este es el electroencefalograma que someteremos posteriormente a estudio. Es a él a quien calculamos los potenciales evocados o, por ejemplo, lo representamos en frecuencia o en tiempo-frecuencia. Es de este electroencefalograma “limpio” del que podemos fiarnos para tomar las medidas que luego podremos incluir en las bases de datos estandarizadas y someter a un estudio estadístico.

## Separación de fuentes en potenciales evocados

La técnica ICA no solo se utiliza para el desartefactado. Obviamente se trata de una aplicación importante y, evidentemente, nos ayuda a obtener una

considerable mejora en la calidad de la información recogida por el EEG. Pero el ICA ha demostrado su utilidad en el estudio de los electroencefalogramas en muchos otros aspectos. En esta sección introducimos una de las numerosas aplicaciones que existen y que, para los protocolos que se describen en el capítulo final del libro, es fundamental. Se trata de la descomposición de los potenciales evocados como superposición de fuentes independientes.

Recordemos que los potenciales evocados se calculan sobre el EEG desartefactado realizando el promedio de un número suficientemente elevado de épocas asociadas al mismo estímulo tipo. Al promediar las respuestas obtenidas entre, por ejemplo, 200 milisegundos antes de la aparición del estímulo y 1.200 milisegundos posteriores al mismo, logramos una buena separación de la señal con la que el encéfalo ha respondido al estímulo, del ruido existente (que es la suma de los ruidos introducidos en el sistema y la actividad eléctrica espontánea).

Una vez se han calculado los correspondientes potenciales evocados, se procede a estudiar las componentes asociadas. Una técnica que se ha mostrado eficaz para la detección de dichas componentes es la aplicación del algoritmo ICA a las curvas ERP. De hecho, cuando procedemos de esta forma, algunos de los potenciales evocados que se conocían clásicamente, como la componente P300, se revelan como la suma de varias componentes nuevas, lo cual muestra que dichos potenciales tienen en realidad una naturaleza más compleja de lo que se creía anteriormente. Lo interesante es que las descomposiciones así calculadas son bastante estables y, por tanto, las nuevas componentes ERP —que llamamos, a partir de ahora componentes ICA del ERP— se pueden incluir en una base de datos que, una vez normalizada, servirá para el diagnóstico y el tratamiento de diversas anomalías clínicas. Obviamente, la curva contra la cual es necesario establecer comparaciones posteriormente, se calcula como un gran promedio a partir de todos los individuos de la base de datos.

El cálculo de las componentes ICA de los potenciales evocados fue introducido en la literatura por S. Makeig y colaboradores a mediados de la

década de 1990 y ha sido utilizado posteriormente por otros investigadores, como Y. Kropotov. En España, lo aplican en sus estudios M Aguilar y colaboradores. Según estos equipos de investigación, la descomposición de las curvas ERP mediante el algoritmo ICA es consistente con el modelo teórico de las componentes ERP, en el que se asume, como vimos en el capítulo primero, que dichas componentes son espacialmente estables, están localizadas en regiones dispersas del encéfalo, y se superponen de forma instantánea. Así, el algoritmo ICA determina qué componentes, temporalmente independientes, se activan para formar las curvas ERP del EEG que hemos grabado. Algunas de las componentes ICA así obtenidas son, en efecto, atribuibles a ciertos procesos neuronales o psicológicos específicos, bien determinados, y, por tanto, tienen el carácter de ser componentes ERP. Estas componentes, una vez identificadas, nos permiten extraer una serie de valores específicos asociados, como son la amplitud de su primer extremo relativo y su latencia (es decir, el tiempo transcurrido desde la emisión del estímulo hasta la aparición de dicho extremo). Estos datos, que provienen de la superposición de numerosos procesos estadísticamente independientes, se puede asumir que responden al comportamiento de una v.a. normal y, por tanto, son susceptibles de ser sometidos a análisis estadístico.

## CAPÍTULO 4

# Neuromarcadores y neuromodulación

### Neuromodulación cerebral: estimulación eléctrica transcraneal y 'neurofeedback'

Los médicos y neuropsicólogos pueden influir de forma cada vez más selectiva en la actividad de las neuronas mediante estímulos eléctricos y magnéticos. Estos métodos han abierto la esperanza de hallar nuevos procedimientos terapéuticos contra el dolor, la depresión, la fibromialgia o el alzheimer, entre otros trastornos. En efecto, se pueden estimular eléctricamente pequeñas redes de neuronas para controlar su actividad y su metaplasticidad, lo cual se ha probado muy útil en numerosos tratamientos clínicos.

La estimulación eléctrica transcraneal se basa en alterar el potencial de membrana de las neuronas mediante una débil corriente eléctrica suministrada por uno o varios electrodos colocados en el cuero cabelludo del sujeto, aplicando campos eléctricos de pocos voltios por centímetro cuadrado. Este procedimiento es no invasivo y produce un leve picor en la zona de aplicación. Funciona disminuyendo el potencial de membrana de las neuronas afectadas cuando aplicamos corriente anódica o, a la inversa, aumentando su potencial de membrana, cuando aplicamos corriente catódica. Este tipo de estimulación transcraneal puede realizarse tanto con corriente continua (ETCC), como con corriente alterna (ETCA), como con ruido aleatorio (tRNS). En este último caso lo que se hace es aplicar desde los electrodos una corriente eléctrica cuya polaridad oscila a altas frecuencias.

Según la posición de los electrodos que se usen en la estimulación, así como el tipo de corriente empleada, se puede forzar o inhibir la actividad cerebral, en



zonas localizadas, de manera temporal. Es, por decirlo en palabras sencillas, como si utilizáramos unos pocos electrodos colocados en el cuero cabelludo, para que ellos realicen las funciones que habitualmente son asumidas por las distintas redes neuronales, bien actuando como “neuronas excitadoras”, bien como “neuronas inhibitoras”, de las neuronas conectadas a ellas por vía postsináptica. Evidentemente, el encéfalo responde bien a este tipo de estímulos, porque en realidad esta respuesta forma parte de su forma natural de actuar, como sabemos. Aunque podría pensarse que los efectos de la electroestimulación son principalmente a nivel cortical, dada la colocación de los electrodos, lo cierto es que este tipo de terapia tiene efectos también a niveles profundos y, de hecho, sus primeras aplicaciones demostradas fueron para el tratamiento de la depresión, que es un fenómeno muy vinculado al sistema límbico.

Otra técnica apropiada para la neuromodulación es el *neurofeedback*. En este caso se persigue que el paciente aprenda a modular por sí mismo ciertos neuromarcadores, como el cociente  $R_{\theta/\beta}$ , cuyo significado neurofisiológico es conocido. El paciente juega con la consola mientras está conectado a unos sensores que registran su EEG y calculan los neuromarcadores que se desea entrenar. El juego solicita que realice ciertos ejercicios, como relajarse o prestar mayor atención, cuya ejecución correcta es detectada al provocar ciertos cambios conocidos en los valores de los neuromarcadores en cuestión. Cuando lo logra, se le premia de algún modo. Obviamente, el paciente está jugando, pero además está aprendiendo a modular su EEG mediante ciertos ejercicios neurocognitivos. De esta forma se pueden entrenar ciertos neuromarcadores para que se aproximen lo máximo posible a sus rangos de normalidad, lo que conlleva una mejoría en la salud.

## Valoración neuropsicológica mediante técnicas de psicometría

En psicología nos interesamos por el estudio científico de la conducta. Para

abordar dicho estudio se miden ciertos parámetros con instrumentos especiales, como cuestionarios, test, evaluaciones de la personalidad, o instrumentos de medida físicos como el EEG. La creación de estos instrumentos de medida, así como su evaluación, su validación, su interpretación y su normalización son el objetivo principal de la psicometría.

El primer instrumento psicométrico que se construyó fue diseñado para medir la inteligencia, a principios del siglo pasado. Otros instrumentos psicométricos son el inventario de personalidad de Minnesota o el modelo de cinco factores para evaluar trastornos de personalidad.

En 1960, Alexander Luria propuso un modelo psicológico basado en funciones básicas cerebrales para explicar los procesos cognitivos subyacentes.

En 1986, D. Norman y T. Shallice propusieron el Modelo de Supervisión de la Atención o SAS. En él se sugiere que los procesos de información en los lóbulos frontales pueden ser modulados por la atención o el control cognitivo y la planificación. Inspirándose en dicho modelo, D. T. Seuss, en 1995, propuso la batería de tests ROBBIA, que divide el sistema atencional en tres subsistemas anatómicos con funciones independientes: energización, monitoreo y entorno de trabajo. La energización se refiere a los procesos de facilitación y potenciación de procesos atencionales, especialmente aquellos encargados de la toma de decisiones y mantenimiento de patrones de respuesta adecuados. El monitoreo proporciona un control de calidad de nuestra conducta, revisando las respuestas a lo largo del tiempo, ajustándose a nuestra conducta. Y el entorno de trabajo se refiere a la formación de criterios, en respuesta a un determinado objetivo, para ejecutar una tarea específica.

Un parámetro importante relacionado con la evaluación psicométrica de la atención es el llamado tiempo de reacción, *TR*. Este se asocia con la capacidad de una persona de realizar lo más rápido posible una tarea concreta. Normalmente, tiempos de reacción elevados se relacionan con actividad metabólica neuronal disminuida en la corteza prefrontal, previa a la presentación de estímulos, o con un incremento en la actividad de la corteza cingular

posterior, el precuneo o el giro temporal medio. La distribución del tiempo de reacción no es gaussiana. Aumenta rápidamente cuando el tiempo de reacción decae y presenta una cola larga cuando este aumenta. Sin embargo, se puede asumir que el logaritmo del tiempo de reacción sigue una distribución gaussiana. Asociados al tiempo de reacción, se calculan los siguientes parámetros:

- Media del tiempo de atención:  $\overline{TR} = \frac{1}{N} \sum_{i=1}^N TR_i$ , donde  $N$  es el número de eventos y  $TR_i$  representa el tiempo de reacción al  $i$ -ésimo estímulo,  $i=1, \dots, N$ .
- Variabilidad del tiempo de reacción:  $VTR = \sqrt{\frac{1}{N-1} \sum_{i=1}^N (TR_i - \overline{TR})^2}$
- Coeficiente de variabilidad:  $CV = VTR / \overline{TR}$ .

Aunque la media y la variabilidad del tiempo de reacción reflejan diferentes mediciones del control cognitivo, ambas están altamente correlacionadas entre sí. Es por ello que se suele tener en cuenta también el coeficiente de variación CV. Hay evidencias suficientes de que la VTR distingue individuos con TDAH de poblaciones sanas, para adolescentes y niños, con magnitud de los efectos grande, y para adultos, con magnitud de los efectos mediana. Por ello la VTR constituye un índice muy importante de estabilidad e inestabilidad del sistema nervioso central de las personas. Sin embargo, pierde especificidad a medida que su valor aumenta.

En resonancia magnética funcional (RMf) los lapsos de atención se relacionan con la actividad generada en la red neuronal por defecto (RND), la cual se mantiene en funcionamiento permanente como actividad basal, emitiendo señales de sincronización que sirven para coordinar la actividad en distintas zonas del encéfalo, asegurándose de este modo que están preparadas para reaccionar de forma conjunta ante los diferentes estímulos que reciben. Se sabe que, cuando el individuo afronta ciertas tareas, como presionar el ratón en una

tarea cognitiva de rendimiento continuo (CPT), ciertas partes de la RND disminuyen su actividad. Es por ello que la activación de las redes defecto anticorrelaciona con las redes positivas y correlaciona positivamente con las fluctuaciones del TR. En otras palabras, la imposibilidad de suprimir la RND, incapacita al sujeto para un rendimiento adecuado en una tarea CPT, o incrementa sus tiempos de respuesta. También es cierto que otras inconsistencias se pueden deber a la práctica, la fatiga o las diferencias de edad.

Los histogramas de la frecuencia de autocorrelación para diferentes parámetros, como los TR, bailan entre 0,01 y 0,1 Hz (fluctuaciones de baja frecuencia). De hecho, se ha comprobado que diferentes parámetros comportamentales, como los tiempos de respuesta, correlacionan con el plano electrofisiológico que medimos en el EEG y con el plano metabólico medido por la RMf.

## **A vueltas con los ritmos del EEG y los potenciales evocados**

Como sabemos, el principal fenómeno EEG es su dinámica oscilatoria. Estas oscilaciones van desde las bajas frecuencias hasta la banda de frecuencias  $\gamma$ . Las oscilaciones del EEG crudo se visualizan mejor en el mundo de las frecuencias, donde el espectro EEG aparece en forma de picos situados en las distintas bandas de frecuencia. Los generadores de estos ritmos están en las diferentes redes neuronales con diferentes tipos de membrana y sinapsis. Además, todas estas neuronas están diseñadas por la evolución para asegurar unas condiciones óptimas de funcionamiento de los procesos de retroalimentación cerebral. Las oscilaciones de los potenciales de acción se producen a través de ciertos canales voltaje asociados, que estabilizan la función de la red neuronal.

Los diferentes patrones del EEG son significativamente diferentes entre individuos sanos. Algunas personas tienen fuerte actividad  $\alpha$  cuando abren los ojos, otros actividad  $\theta$  frontal media, y otros no muestran ni actividad  $\alpha$  ni  $\theta$ . La amplitud de la actividad EEG en una persona depende de muchos factores, además de su genética, neurofisiología, y las propiedades físicas o anatómicas de

los tejidos circundantes al cerebro (piel, hueso, duramadre y piamadre). Estos parámetros cambian de un sujeto a otro, y son desconocidos. Para compensar estas variaciones se calcula la energía relativa del EEG.

De acuerdo con las especificaciones vigentes desde hace aproximadamente 50 años, el espectro de banda en el que se asume que oscila la actividad del encéfalo va desde los 0,5 a los 50 Hz. Se escogieron entonces filtros paso alto para eliminar potenciales eléctricos generados por fuentes externas al cerebro, como los artefactos musculares o la red eléctrica. Por otra parte, tanto los potenciales eléctricos de altas como de bajas frecuencias han sido investigados en pocos laboratorios. Parece que ambos están relacionados con oscilaciones metabólicas como el flujo sanguíneo o el oxígeno extracelular. El interés actual en el estudio de la actividad eléctrica en estas bandas es mayor a partir del descubrimiento de que las fluctuaciones de baja frecuencia tienen lugar en las señales BOLD de la RMf y, también, porque, a partir de la comercialización de amplificadores EEG de banda EEG completa, se aplican con éxito en el tratamiento de trastornos mentales, mediante *neurofeedback*.

Aunque los potenciales evocados se miden con los mismos amplificadores con los que se mide la actividad EEG espontánea, desde el punto de vista funcional hay una diferencia fundamental entre estas dos mediciones, pues los potenciales evocados reflejan estados en el procesamiento de información en respuesta a un estímulo. Estos estados pueden incluir los patrones del flujo de activación o inhibición en una red de neuronas sensoriales jerárquicamente organizadas, inducidos por estímulos sensoriales, o los patrones de flujo de activación o inhibición de redes del sistema ejecutivo que generan una acción final. Se muestran por tanto dos tipos de actividad contrapuesta: actividad sensorial y actividad motora. En los sistemas sensoriales predomina la primera y, en los sistemas de control cognitivo o motor, predomina la segunda. Los sistemas sensoriales extraen información desde la entrada y seleccionan la parte de información según mecanismos de atención. El sistema emocional gestiona respuestas fisiológicas ante estímulos emocionales. El sistema de la memoria

añade experiencias previas a la situación actual y consolida la memoria para los eventos futuros. El sistema cognitivo planifica e inicia acciones relevantes, suprime acciones irrelevantes y monitorea conductas para un control de calidad. Estos sistemas se distribuyen en diferentes partes del cerebro y se activan en diferentes intervalos de tiempo.

La amplitud de los potenciales evocados es muy pequeña, por lo que es imposible separarla a simple vista de la actividad espontánea de fondo. En otras palabras, la relación señal/ruido para un estímulo determinado es muy baja. Para incrementar esta relación, se promedian gran cantidad de fragmentos del EEG, para un elevado número de estímulos o acontecimientos. Por ejemplo en un test de paradigma GO/NOGO hay diferentes categorías de estímulos (como ignorar, GO y NOGO) que se presentan, aleatoriamente, muchas veces. Si seleccionamos la condición ignorar, en la que el estímulo visual se presenta pero no se requiere respuesta del sujeto, el EEG completo se divide en pequeños trozos, que llamamos épocas, correspondientes a cada estímulo de este tipo, que consisten en una parte preestímulo y una parte postestímulo. Los potenciales evocados asociados a dicha condición se obtienen entonces como resultado de promediar todos estos fragmentos del EEG. Normalmente, cuando promediamos 200 estímulos para una persona en particular, las fluctuaciones EEG que preceden al estímulo se cancelan unas a otras, aproximándose a 0 en dicha zona. La relación señal/ruido se define no solo en base a la amplitud de la onda medida, sino también a la amplitud de la actividad de fondo. Es por ello que medir la onda N1 o C1 generada en la corteza visual es más difícil que la medición de componentes tardíos, como la onda P300.

Hay varios tipos de paradigmas para los potenciales evocados. Concretamente, se separan en varias categorías que dependen del sistema a estudiar, y son: tareas sensoriales, motoras, atencionales, emocionales, de memoria y de control cognitivo.

Los componentes ERP sensoriales son modulados por estímulos visuales, auditivos o somatosensoriales, con características físicas con una modalidad

(posición, color, orientación, etc.), por categorías de estímulos (objetos animados, objetos inanimados, caras, herramientas, etc.), por la frecuencia de presentación (por ejemplo intervalos de 100 a 2.000 ms interestímulo), o por la intensidad del estímulo (por ejemplo volumen del sonido). Los paradigmas escogidos suelen ser test de tareas infrecuentes en modalidades auditiva o visual, tareas de escucha dicótica, etc.

Para estudiar el sistema motor, los potenciales evocados se relacionan con acciones motoras en vez de estímulos sensoriales. La acción puede autoiniciarse o ser contingente con algún estímulo sensorial. Por ejemplo, según Posner, para estudiar la atención hace falta un paradigma con al menos tres estímulos.

Las tareas de control cognitivo imitan situaciones de comportamiento cuando un modelo preponderante de ejecución se forma por la configuración de tareas y domina durante la tarea, pero en algunos casos se ha violado, de manera que el sujeto tiene que anular la acción preparada con una acción diferente o, simplemente, suprimirla.

## **La base de datos SolidBrain**

Cualquier parámetro, cuando se mide en una población dada, varía de una persona a otra. Esto sucede, evidentemente, con todos los datos que se puedan obtener a partir de la grabación de un EEG y su estudio posterior. En particular, esto se produce con todos los neuromarcadores que podemos definir por esta vía. Cuando un paciente acude a consulta para que se le calculen estos neuromarcadores, es fundamental que se obtengan siguiendo un protocolo muy estricto, que coloque en condiciones de igualdad a todos los pacientes.

El EEG se graba para identificar, en cada paciente, el valor de los distintos neuromarcadores al uso, que en el caso que nos ocupa son las latencias y amplitudes de las diferentes componentes ICA de los potenciales evocados calculados, siguiendo los distintos paradigmas que existen (que nos informan sobre aspectos relacionados con los sistemas sensorial, emocional, ejecutivo y de memoria), así como las diferentes medidas espectrales (en frecuencia, y en

tiempo-frecuencia), tanto de los potenciales evocados como de la actividad EEG sin promediar.

Todos estos neuromarcadores son variables aleatorias y, para poder realizar diagnósticos a partir de ellos, es necesario disponer de una buena base de datos normalizada contra la cual se han de contrastar los datos obtenidos paciente a paciente. Sobre la naturaleza de las variables aleatorias implicadas, hay que decir que se sabe que las amplitudes y las latencias de los potenciales evocados se distribuyen normalmente. Sin embargo, los neuromarcadores obtenidos en el dominio de la frecuencia (o en tiempo-frecuencia) tanto para actividad EEG espontánea como para potenciales evocados, no son variables aleatorias normales, por lo que antes de ser estudiados es necesario someterlos a una conocida transformación matemática, basada en un cambio de escala logarítmico, que, se sabe, tiene el efecto de reconvertir dichas variables en variables aleatorias cuyo comportamiento es aproximadamente normal. Este proceso se denomina “normalización”. Evidentemente, construir una buena base de datos es un proceso lento y muy costoso. Pero es gracias a la existencia de este tipo de bases de datos que podemos esperar la llegada de nuevos logros científicos, que beneficiarán no solo a la comprensión de numerosos fenómenos sino también al sistema de salud en general.

Desde hace más de una década, M. Aguilar y colaboradores han analizado los EEG y los potenciales evocados de miles de sujetos con diferentes categorías diagnósticas (muchas veces sin una causa clara de enfermedad). Para comparar el grupo de pacientes con sujetos sanos se ha construido, entre 2011 y 2013, la base de datos SolidBrain. Las grabaciones fueron realizadas en España y México.

Aunque se han analizado por encima de 20.000 sujetos, en la base de datos solo hay incluidos 8.000, que están completamente normalizados, de los cuales 2.000 son sujetos sanos o controles y 6.000 son pacientes. Los datos se han tomado tanto en estado de reposo (en cuyo caso se mide solamente la actividad eléctrica espontánea) como estado de respuesta a tareas con estímulos acústicos,



visuales, etc., siguiendo los diferentes paradigmas existentes para el estudio de potenciales evocados: GO, NOGO, MMN. Los sujetos sanos, así como sus parientes más cercanos (de primer y segundo grado), carecen de historia de problemas neurológicos o psiquiátricos. La recolección de datos normalizados se inició en Elche por F. Vargas-Torcal, quien se inspiró en el trabajo realizado por Aguilar en la Universidad de Murcia. La metodología empleada había sido usada con anterioridad en la Academia de Ciencias rusa, por Y. Kropotov. Para mejorar la relación señal/ruido de los potenciales evocados se impone un mínimo de 100 de muestras por cada categoría de estímulos a analizar, en cada individuo.

Esta base de datos se ha realizado con metodología desarrollada por la empresa SolidBrain, en colaboración con la Universidad de Murcia y la Academia de Ciencias Rusa, aunque el código fuente es propiedad de SolidBrain, que tiene su sede en Teruel.

Las claves que justifican la creación de una base de este tipo, para su uso posterior en estudios clínicos, son las siguientes:

- Los potenciales EEG recogidos en cuero cabelludo son el resultado de la suma de varias fuentes con diferentes localizaciones (redes neuronales en una coordenada concreta entre millones de neuronas interconectadas entre sí) y diferentes propiedades funcionales.
- Se ha logrado asociar cada componente de los potenciales evocados extraídos de la tarea GO/NOGO por ICA, con índices neuropsicológicos descritos según los modelos clásicos de la neuropsicología como el propuesto por Luria. Dicho en palabras más sencillas: cada una de las componentes ICA de los potenciales evocados tiene un significado funcional en función de su amplitud y su latencia, y se relaciona con cada uno de los sistemas del encéfalo: sensorial, memoria, emocional o ejecutivo.
- Las técnicas de neuromodulación, como la estimulación eléctrica transcraneal y el *neurofeedback*, ayudan a modificar estas curvas hacia

arriba o abajo. Además, si no se aplica ningún tipo de terapia, estas curvas se mantienen estables en el tiempo. En particular, mediante estimulación eléctrica transcraneal podemos subir o bajar la curva ICA calculada, según apliquemos una polaridad u otra: la corriente anódica aumenta el voltaje y la catódica lo disminuye.

El uso de neuromarcadores obtenidos del EEG y de los potenciales evocados, comparando los resultados calculados en cada paciente con los controles sanos de una base de datos normalizada, sirve para definir protocolos individualizados de tratamiento, tanto mediante *neurofeedback* como con estimulación eléctrica transcraneal, que ayudan a mejorar la calidad de vida de los pacientes. Los resultados han sido muy satisfactorios, y los efectos secundarios, reducidos.

En España existen varias clínicas donde se siguen estos protocolos, basados en la grabación del EEG y la comparación con la base de datos SolidBrain.

## Diagnóstico basado en neuromarcadores

Los trastornos mentales son un problema global y representan uno de los mayores retos para los principales sistemas de salud a nivel mundial. En el mundo hay en torno a 500 millones de personas que sufren algún trastorno mental, y en la Unión Europea, los trastornos mentales se extienden como una de las principales causas de carga económica y enfermedad. Lo que empeora la situación es que se espera que la prevalencia de los trastornos mentales no pare de crecer, por varias razones, incluyendo el envejecimiento de la población, el empeoramiento de la situación económica de muchas familias, los problemas de seguridad laboral reducida, el aumento de las horas de trabajo y el estrés laboral.

En el pasado los trastornos mentales se definían en ausencia de lesiones orgánicas. Aquellos trastornos asociados a lesiones cerebrales pasaban automáticamente a considerarse trastornos neurológicos. Como habitualmente en los siglos XIX y XX no se encontraba daño estructural, se diagnosticaba a los pacientes siguiendo criterios clínicos, encapsulados en forma de categorías

diagnósticas compuestas por una descripción de síntomas y signos. Estas categorías diagnósticas, definidas en el *DSM-5*, suponen una autoridad a nivel internacional, pues en la práctica clínica determinan los diagnósticos y, para las compañías de seguros, suponen la estrategia de pago o indemnización correspondiente a cada diagnóstico.

Evidentemente, esta metodología diagnóstica queda en entredicho cuando la comparamos con la práctica usual en otras disciplinas médicas, como la cardiología o la oncología, donde los diagnósticos se basan en pruebas objetivas de laboratorio, en vez de consensos vagos de agrupaciones de síntomas. Además, sabemos por estas disciplinas médicas que los síntomas no afianzan, por sí mismos, un tratamiento médico, y menos aún si es intervencionista. Solo cuando se tienen en cuenta la clínica y los test de laboratorio, se pueden tomar medidas oportunas.

Entonces surge la pregunta: ¿por qué deberíamos usar grabaciones de señales eléctricas para el diagnóstico de las enfermedades del corazón, pero no usarlas en psiquiatría? En realidad, esto es algo que está en proceso de cambio, gracias a las numerosas publicaciones que existen avalando la utilidad de la electroencefalografía cuantitativa, así como otros métodos de neuroimagen, cada vez más en uso. Sin embargo, resulta necesario explicar aquí algunas de las razones que pudieran revelar la prevalencia de esta tendencia hasta fechas recientes.

En primer lugar, una diferencia importante cuando comparamos con la cardiología es la complejidad: el cerebro es mucho más complejo que el corazón. La fisiología del corazón proporciona un número relativamente limitado de parámetros del electrocardiograma (ECG) que se pueden utilizar para el diagnóstico. Los parámetros del ECG incluyen solo cinco formas de onda, que son identificadas fácilmente en cada sujeto y tienen un significado funcional claro en cada ciclo cardíaco. En contraste el número de parámetros proporcionado por un EEG multicanal en estado de reposo es enorme, incluso si utilizamos el sistema internacional 10-20, que no requiere de un número elevado

de sensores (actualmente se pueden grabar los EEG con hasta 300 sensores). Si los potenciales evocados se incluyen en el análisis, la cantidad de información disponible aumenta enormemente.

En segundo lugar, las formas de onda del ECG no muestran grandes diferencias entre individuos y tienen un valor diagnóstico claro. En contraste, el EEG y, sobre todo, los parámetros ERP, muestran gran variación interindividual. Encima los parámetros del EEG/ERP son muy sensibles a las fluctuaciones en el estado de los sujetos. Parece ser que la población sana no es homogénea en términos de funcionamiento del cerebro. Si tenemos en cuenta la heterogeneidad de una categoría diagnóstica para una determinada enfermedad mental, las dificultades en la separación de los pacientes con un diagnóstico determinado, a partir de los controles sanos, mediante un biomarcador, se hace evidente. Sin embargo, en defensa del EEG, diremos que, afortunadamente, las curvas obtenidas a través de la técnica de gran promedio son estables y, por tanto, una vez calculadas para una base de datos estandarizada, se pueden utilizar para cuestiones clínicas.

En tercer lugar, el registro y los procedimientos de medición de ECG están completamente estandarizados. En contraste, en el caso del EEG ha sido estandarizada solamente la colocación de electrodos (sistema internacional 10-20). Los otros parámetros, como las condiciones de grabación de los potenciales evocados, no están aún estandarizados. Como consecuencia, la recogida de bases de datos en la práctica clínica es bastante limitada en el campo del EEG y prácticamente ausente en el campo de los potenciales evocados.

En cuarto lugar, el corazón es un órgano relativamente simple (básicamente, es una bomba electromecánica), mientras que el cerebro es el sistema más sofisticado del mundo, que está destinado no solo a autorregularse a sí mismo, sino también a procesar la información externa e interna, consciente e inconsciente. Un cardiólogo, en su práctica clínica, se basa en una teoría del corazón bien establecida. No existe, sin embargo, una teoría definitiva para el funcionamiento del cerebro. Razones similares pueden explicar por qué otros

marcadores neurofuncionales tales como los que se obtienen a través de MEG, fMRI, o PET no se han estandarizado aún en la práctica clínica.

Las categorías de diagnóstico de la enfermedad mental que se describen en el *DSM-5* no están organizadas de acuerdo con disfunciones de los sistemas cerebrales, sino por sintomatologías. Sin embargo, como muestra la neuropsicología, síntomas similares pueden ser inducidos por un daño local en diferentes partes de un circuito neuronal, ampliamente distribuidos por el cerebro. Esta afirmación se apoya en estudios recientes de la actividad funcional del cerebro (basados en potenciales evocados, fMRI, PET, etc.) registrada en varios trastornos psiquiátricos, en comparación con la actividad cerebral funcional en controles sanos. Cada vez más estudios demuestran claramente que los patrones de activación cerebral pueden tener un valor diagnóstico, tal como lo tiene la formación de imágenes cardíacas durante una prueba de esfuerzo, las cuales se utilizan ahora, por ejemplo, para diagnosticar la enfermedad arterial coronaria.

El *DSM-5* se publicó el 18 de mayo de 2013. Fue criticado por diversas autoridades por su falta de apoyo empírico y por la influencia de la industria farmacéutica. En psiquiatría se estaba esperando el reconocimiento de distintos biomarcadores de trastornos mentales, tema en el que se había trabajado intensamente durante las últimas décadas. Sin embargo, estos no han sido finalmente recogidos, o se han considerado solo muy parcialmente, en la versión final del texto. Así pues, la publicación definitiva del *DSM-5*, ha evocado sentimientos de decepción. No obstante, el 15 de julio de 2013, la FDA permitió la comercialización del primer dispositivo médico basado en EEG para ayudar a evaluar el TDAH en niños y adolescentes de 6 a 17 años. El dispositivo que ayuda a la evaluación de registros EEG se denomina NEBA y recoge señales EEG durante 15-20 minutos, calculando el cociente  $R_{\theta/\beta}$ , que ha demostrado estar aumentado en niños y adolescentes con TDAH, cuando son comparados con controles sanos. Esto se considera un síntoma de que, en efecto, es cuestión de

tiempo que, en futuras versiones del *DSM*, se tengan en consideración nuevos neuromarcadores.

Todavía no podemos establecer una tabla que asocie de manera causal el comportamiento de los distintos neuromarcadores hallados a partir del EEG, con categorías diagnósticas concretas, aunque muchos pensamos que esto podría producirse en un futuro más o menos cercano. Mientras tanto, lo que sí está claro es que individuos diagnosticados de diversos trastornos mentales presentan valores en ciertos neuromarcadores, que son predecibles, al igual que sucede con los individuos sanos. Además, estos neuromarcadores pueden modularse mediante electroestimulación y *neurofeedback*, y cuando logramos que se acerquen a zonas de normalidad, el estado de salud de los pacientes mejora, y lo hace en aspectos que son también predecibles, como el aumento de su atención, su sociabilidad o su mejoría en aspectos motrices o cognitivos.

### **Pero ¿qué es un neuromarcador?**

Intuitivamente un biomarcador (o marcador biológico) es una característica que puede ser medida y evaluada objetivamente, y que sirve como indicador de procesos biológicos normales, procesos patógenos o respuestas farmacológicas a una intervención terapéutica. De acuerdo con el tipo de información que proporcionan, los biomarcadores para trastornos del sistema nervioso central, también llamados neuromarcadores, pueden ser clasificados como clínicos, basados en neuroimagen, bioquímicos, genéticos o marcadores proteómicos. Las expectativas para el desarrollo de nuevos neuromarcadores en cada una de estas categorías, y, en particular, basados en la medición del EEG, son altas, y se espera que causen una mejora significativa en el diagnóstico y la prevención de enfermedades neurológicas y psiquiátricas.

La definición de neuromarcadores concretos para los distintos trastornos mentales es una tarea importante y difícil, en la que se trabaja de forma activa desde hace tiempo. Como ejemplo de los requisitos que se requieren para que una cierta medida sea admitida como neuromarcador para alguna patología,

mencionamos el caso específico de los que se usan en el diagnóstico del TDAH. Ellos fueron delimitados en un informe de consenso, alcanzado en 2012, por la Federación Mundial del TDAH. Según dicho informe, un neuromarcador ideal para el TDAH debe tener:

- una sensibilidad diagnóstica superior al 80% para la detección de TDAH,
- una especificidad diagnóstica superior al 80% para distinguir el TDAH de otros trastornos con síntomas similares.

Además,

- debe ser fiable, reproducible, barato de medir, no invasivo y fácil de calcular,
- debe ser confirmado por al menos dos estudios independientes, llevados a cabo por investigadores cualificados, con los resultados publicados en revistas de prestigio.

Los neuromarcadores se deben expresar numéricamente de forma individualizada, para cada sujeto. Sabemos, además, que cualquier medición se realiza con un margen de error. Por tanto, una cuestión fundamental es averiguar qué mecanismos podemos utilizar para estimar y/o reducir estos errores.

Fue el psicólogo inglés C. E. Spearman quien, a principios del siglo XX, y preocupado por la elaboración de algún tipo de medida general para la inteligencia, abordó por primera vez el estudio de esta cuestión concreta, que es fundamental para la psicometría. Creó los fundamentos de la teoría clásica de los tests. Dicha teoría asume que, fijado un test de medición psicométrica, cada persona tiene una puntuación real (o verdadera)  $T$  asociada a dicho test, que se obtiene si no hay errores en la medición. La puntuación verdadera de un sujeto se obtiene numéricamente como la esperanza matemática de la variable aleatoria que resulta de aplicar el test a dicha persona. Obviamente, este valor no se puede

conocer por la mera aplicación del test un número determinado de veces. Por el contrario, cada vez que aplicamos un test o prueba psicométrica (y esto incluye, evidentemente, la grabación del EEG y la toma posterior de medidas asociadas) solo obtendremos una observación, que es susceptible de ser medida, y que se descompone como  $X=T+e$ , donde  $X$  es la variable observada,  $T$  el valor verdadero del test, y  $e$  es el error cometido. Por definición,  $T=E(X)$ . El error  $e$  se descompone como la suma  $e = E_r + E_s$  de un error de naturaleza aleatoria,  $E_r$ , y un error  $E_s$  llamado sistemático porque permanece dentro de un mismo umbral cuando el experimento se repite (por ejemplo, una fecha tomada mal o una dirección mal dada, o la inclusión de información irrelevante para nuestros objetivos de medida, etc.). El error aleatorio se produce como efecto de la variación aleatoria del mismo sujeto a lo largo de los diferentes tests. Las fuentes de este tipo de error son por ejemplo las variaciones del humor del sujeto experimental, su fatiga, su estrés, el sesgo de la motivación del sujeto, las trampas, las variaciones del ruido del entorno, los cambios de temperatura, la iluminación, la comodidad de su posición mientras realiza una prueba, cómo se explican las instrucciones del test o cómo se administra, cómo se puntúan los aciertos y/o errores, etc. En el caso particular de que estemos evaluando la actividad cerebral recogida por un EEG, el error aleatorio podría tener su origen en variaciones de la actividad espontánea para un determinado estado cerebral, como también en errores del instrumento de medida. Las variaciones espontáneas del parámetro pueden ser disminuidas por técnicas de promediado —por eso se calculan los potenciales evocados—. De hecho, cuando promediamos, la relación señal-ruido, SNR, que es el parámetro utilizado por defecto en este contexto para determinar la calidad de una medición dada, y que se define como el cociente de la energía de la señal (sin errores) por la energía del ruido recogido, se incrementa multiplicando por la raíz cuadrada del número de estímulos promediados. Por supuesto, a mayor SNR, mayor calidad de los datos. Finalmente, los errores del instrumento de medida vienen determinados



por los fabricantes y no pueden ser disminuidos.

Algunos conceptos importantes asociados a los tests diagnósticos, y que se utilizan para establecer la calidad de los neuromarcadores, son la repetibilidad, la reproducibilidad, la confianza y la validez. Resumimos muy brevemente lo que significan estos conceptos: La repetibilidad de las mediciones se refiere a la variación en las mediciones repetidas realizadas sobre el mismo sujeto en idénticas condiciones y durante un periodo corto de tiempo (varias horas o días). La variabilidad en las mediciones realizadas en estas condiciones se puede atribuir solamente a los errores de medición  $E_r$ . La reproducibilidad se refiere a la variación en las mediciones realizadas sobre un sujeto en condiciones cambiantes. Las condiciones cambiantes pueden deberse a diferentes métodos o instrumentos de medición que se utilizan, a que las mediciones se realizan por diferentes observadores, o también a que las mediciones se realizan espaciadas por un largo periodo de tiempo (de varias semanas a varios meses). La validez se refiere a la fiabilidad con la que una prueba mide lo que se supone que debe medir. El error sistemático  $E_s$  afecta a la validez, pero no a la repetibilidad o la reproducibilidad del test. La confianza del test se calcula como el cociente de la varianza del valor verdadero del mismo por la varianza del valor observado (téngase en cuenta que un mismo sujeto puede cambiar con el tiempo, con terapias, etc., por lo que el valor verdadero del test aplicado a dicho individuo es,

también, una variable aleatoria):  $\rho_{XT}^2 = \frac{\sigma_T^2}{\sigma_X^2}$ . Como  $X=T+e$  y  $T, e$  son independientes,  $\sigma_X^2 = \sigma_T^2 + \sigma_e^2$  y, sustituyendo en la ecuación anterior,

obtenemos la nueva fórmula para la confianza del test:  $\rho_{XT}^2 = 1 - \frac{\sigma_e^2}{\sigma_X^2}$ . A partir de esta ecuación se observa que la confianza de los resultados de las pruebas es siempre menor que 1, se hace mayor cuando la varianza del error  $\sigma_e^2$  disminuye y, además, la heterogeneidad de los sujetos en la población, medida por la

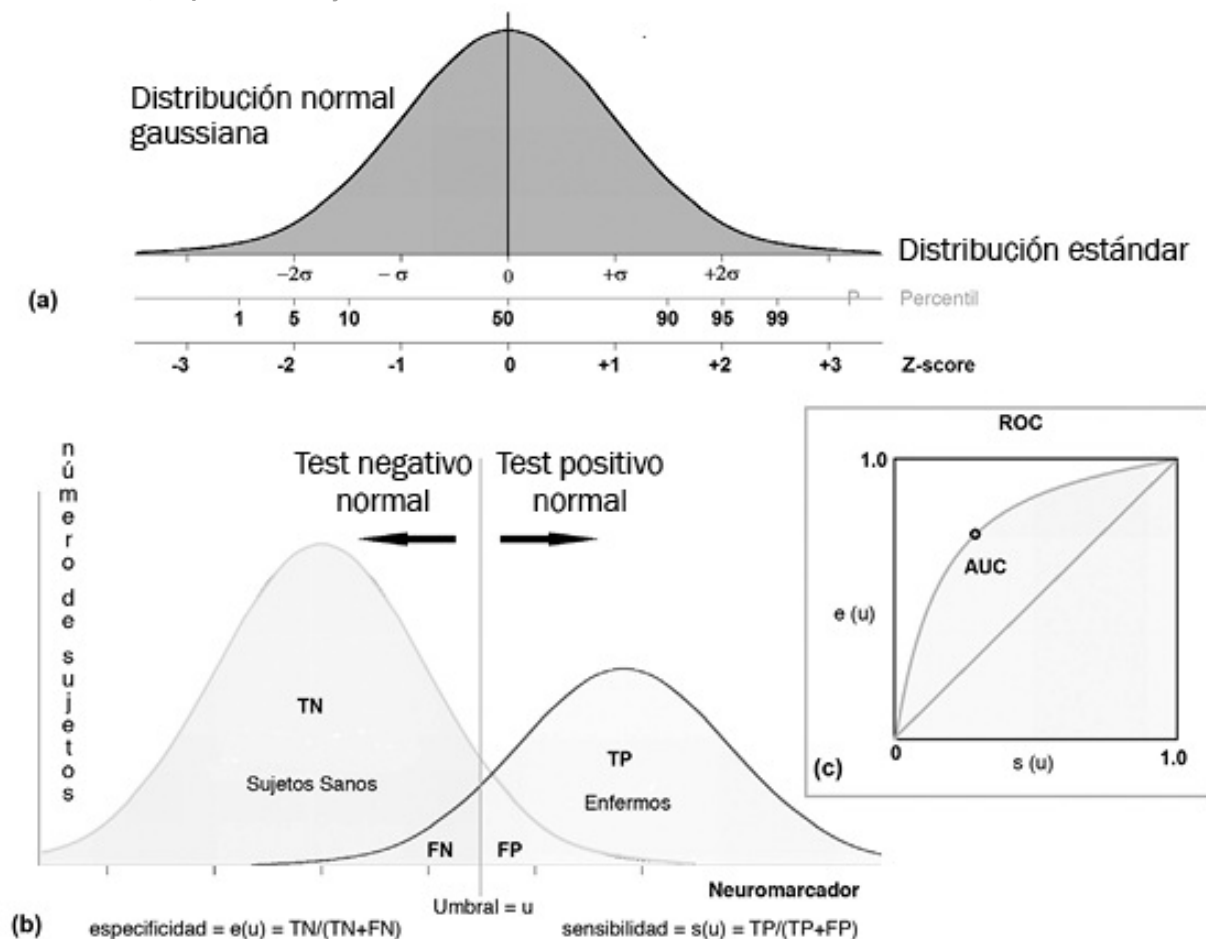
varianza de la variable observada  $\sigma_x^2$ , afecta al valor de confianza en el sentido de que cuanto mayor es la heterogeneidad, más alta es la confianza. Se puede demostrar, además, que la raíz cuadrada de la confianza, el llamado el índice de confianza  $\rho_{XT}$ , es igual a la correlación entre las puntuaciones observadas y las puntuaciones verdaderas en la población general. En la práctica, la correlación entre las puntuaciones observadas  $X$ , medidas en dos momentos diferentes en una población seleccionada de sujetos, se utiliza para estimar la confianza del valor  $X$  medido. Esta correlación mide la confianza del test-retest. Otra medida de la confianza que se puede introducir de forma adicional, para evaluar los distintos neuromarcadores, se obtiene dividiendo al azar toda la muestra poblacional en dos grupos iguales y calculando el coeficiente de correlación de Pearson entre ambos conjuntos.

Aunque es necesario disponer de una estimación de la confianza para cada neuromarcador, ya que una confianza elevada nos garantiza la medición apropiada del mismo, este valor por sí solo no es suficiente garantía de su bondad. Para que una medida sea fiable, también debe ser válida. Como se mencionó anteriormente, la validez es el grado en que una medida está bien fundada y se corresponde con precisión con el mundo real. En el caso de los neuromarcadores apoyados en la grabación del EEG, los estudios sobre su validez se apoyan en conocimientos de neurofisiología, neuropsicología, etc., así como en la repetición de experimentos. Además, se sabe que son confiables (tienen un índice de confianza  $\rho_{XT}$  elevado), repetibles y reproducibles.

Introducimos ahora los conceptos de sensibilidad y especificidad, ya que estos aparecen como parte fundamental en la definición de neuromarcador ideal para el diagnóstico del TDAH que expusimos anteriormente. Supongamos que disponemos que un neuromarcador para detectar qué personas padecen una determinada enfermedad. Cada persona a la que medimos el neuromarcador puede tener o no tener la enfermedad. El resultado de la prueba, consistente en

medir el neuromarcador y comprobar si este supera (o no) un cierto umbral  $u$ , puede ser positivo (predice que la persona tiene la enfermedad) o negativo (predice que la persona no tiene la enfermedad). Los resultados de la prueba para cada sujeto pueden o no coincidir con la situación real del sujeto. La sensibilidad  $s(u)$  del neuromarcador, fijado un umbral  $u$  dado, se define, entonces, como la probabilidad de que la prueba sea positiva dado que el sujeto está enfermo. La especificidad  $e(u)$  es la probabilidad de la prueba proporcione un resultado negativo dado que el sujeto está sano.

FIGURA 17  
Sensibilidad, especificidad y curva ROC.



Para cualquier prueba, existe un compromiso entre sensibilidad y especificidad. Esta compensación se puede representar gráficamente a través de

la curva de rendimiento diagnóstico (ROC), cuyo gráfico resulta de dibujar, para cada umbral de decisión  $u$ , el par ordenado  $(s(u), 1 - e(u))$ . Dada esta curva, su punto más cercano al punto  $(0,1)$ , la esquina superior izquierda del cuadrado en el que está inscrita, determina el mejor umbral de decisión posible asociado al neuromarcador. Además, las curvas ROC se pueden utilizar para elegir entre dos pruebas diagnósticas distintas (o también entre dos neuromarcadores). La elección se realiza mediante la comparación del área bajo la curva (AUC) de ambas pruebas. Dicho área posee un valor comprendido entre 0,5 y 1, donde 1 representa un valor diagnóstico perfecto y 0,5 es una prueba sin capacidad discriminatoria diagnóstica. Se considera que el neuromarcador es muy bueno si el área bajo la curva de rendimiento diagnóstico asociada supera el valor 0,9, y se considera simplemente bueno cuando este valor supera a 0,75.

La década de 1990 fue conocida como “década del cerebro”. Entonces se formularon nuevos conceptos de funcionamiento del encéfalo, basados en los datos acumulados gracias a la implantación de técnicas de neuroimagen. También se desarrollaron nuevos enfoques metodológicos, tales como técnicas de separación ciega de fuentes para la detección de las fuentes locales de actividad cerebral. Los primeros quince años de este siglo podrían ser categorizados como la “década de los descubrimientos”, pues en este periodo de tiempo se han identificado numerosos circuitos neuronales de funcionamiento normal y anormal del cerebro. Ahora es el momento en que los nuevos métodos de tratamiento, no necesariamente farmacológicos, como la estimulación eléctrica transcraneal, la estimulación magnética transcraneal y otras formas de neuromodulación, se han probado en la práctica clínica. La década de los descubrimientos va a ser seguida, con toda probabilidad, por la “década de la translación” que se centrará en la identificación de neuromarcadores asociados a cada uno de los principales trastornos mentales, para facilitar la detección temprana y la prevención, así como la atención personalizada para cada paciente. Además, pensamos que la detección temprana de enfermedades mentales basada

en el uso de neuromarcadores facilitará el desarrollo de intervenciones preventivas.

# Bibliografía

- AMIDROR, I. (2013): *Mastering the Discrete Fourier Transform in One, Two or Several Dimensions*, Computational and Imaging Vision, 43, Springer.
- ARÉCHIGA, H. (2001): “El universo interior”, *La ciencia para todos*, 182, Fondo de Cultura Económica.
- BELL, A. J. y SEJNOWSKI, T. J. (1995): “An information maximization approach to blind separation and blind deconvolution”, *Neural. Comput.*, 7, pp. 1129-1159.
- (1999): “Independent component analysis using an extended infomax algorithm for mixed sub-Gaussian and super-Gaussian sources”, *Neural. Comput.*, 11, pp. 417-441.
- BRAZIER, M. A. B. (1976): *Actividad eléctrica del sistema nervioso*, Ediciones Literarias y Científicas.
- BUZSÁKI, G. (2006): *Rhythms of the Brain*, Oxford University Press.
- BUZSÁKI, G. et al. (2012): “The origin of extracellular fields and currents —EEG, ECoG, LFP and spikes”, *Nature reviews. Neuroscience*, 13(6), pp. 407-420.
- COHEN, M. X. (2016): *Cycles in the mind. How brain rhythms control perception and action*, Sinc(X) Press.
- (2014): *Analyzing time series data. Theory and Practice*, The MIT Press.
- DEMOS, J. N. (2005): *Getting started with Neurofeedback*, Ed. Norton & Co.
- DUFFY, F. H. (1986): *Tophographic mapping of brain electrical activity*, Londres, Butterworths.
- GORDON, E. (2007): “Integrating genomics and neuromarkers for the era of brain-related personalized medicine”, *Personalized Medicine*, 4 (2), pp. 201-217.
- GORDON, E. et al. (2005): “Integrative neuroscience: the role of a standardized database”, *Clinical EEG and Neuroscience*, 36 (2), pp. 64-75.
- HAUSER, M. W. (1991): “Principles of oversampling A/D conversion”, *J. Audio Eng. Soc.*, 39 (1-2), pp. 3-26.
- HODGKIN, A. L. y HUXLEY, A. F. (1939): “Action potentials recorded from inside a nerve fibre”, *Nature*, 144, pp. 710-711.
- (1952): “A quantitative description of membrane current and its application to conduction and excitation in nerve”, *J. Physiology*, 117 (4), pp. 500-544.
- HYVÄRINEN, A. et al. (2001): *Independent Component Analysis*, John-Wiley & Sons, Inc.
- KROPOTOV, Y. et al. (2011): “In Search of New Protocols of Neurofeedback: Independent Components of Event-Related Potentials”, *Journal of Neurotherapy*, 15 (2), pp. 151-159.
- LOO, S. K. et al. (2016): “Research Review: Use of EEG biomarkers in child psychiatry research —current state and future directions”, *Journal of Child Psychology and Psychiatry*, 57 (1), pp. 4-17.
- LUCK, S. J. y KAPPENMAN, E. S. (2013): *The Oxford Handbook of Event-Related Potential Components*, Oxford University Press.
- LÜKE, H. D. (1999): “The origins of the sampling theory”, *IEEE Comm Mag.*, 37 (4), pp. 106-108.
- LURIA, A. R. (1980): *Higher cortical functions in man*, Nueva York, Basic Books.

- MAKEIG, S. *et al.* (1997): “Blind separation of event-related brain responses into independent components”, *Proceedings of the National Academy of Sciences of the United States of America*, 94, pp. 10979-10984.
- MONASTRA, V. J. *et al.* (1999): “Assessing attention deficit hyperactivity disorder via quantitative electroencephalography: an initial validation study”, *Neuropsychology*, 13, pp. 424-433.
- NITSCHKE, M. A. *et al.* (2002): “Modulation kortikaler Erregbarkeit beim Menschen durch transkranielle Gleichstromstimulation”, *Der Nervenarzt*, 73 (4), pp. 332-335.
- (2000): “Excitability changes induced in the human motor cortex by weak transcranial direct current stimulation”, *The Journal of Physiology*, 527 (3), pp. 633-639.
- NORMAN, D. A. y SHALLICE, T. (1986): “Attention to action: Willed and automatic control of behaviour”, en R. J. Davidson, G. E. Schwartz y D. Shapiro (eds.), *Consciousness and Self-Regulation: Advances in Research and Theory*, vol. 4, Plenum Press, pp. 1-18.
- ONTON, J. y MAKEIG, S. (2006): “Information-based modeling of event-related brain dynamics”, *Prog. Brain Res.*, 159, pp. 99-120.
- OPITZ, A. *et al.* (2015): “Determinants of the electric field during transcranial direct current stimulation”, *Neuroimage*, 109, pp. 140-150.
- PERRIN, F. *et al.* (1989): “Spherical splines for scalp potential and current density mapping”, *Electroencephalography and Clinical Neurophysiology*, 72 (2), pp. 184-187.
- SEUNG, S. (2012): *Conectoma*, Barcelona, RBA.
- SHALLICE, T. (1988): *From neuropsychology to mental structure*, Cambridge, UK, Cambridge University Press.
- STONE, J. V. (2004): *Independent Component Analysis. A tutorial introduction*, The M.I.T. Press.
- VANHATALO, S. *et al.* (2005): “Full-band EEG (fbEEG): a new standard for clinical electroencephalography”, *Clinical EEG and Neuroscience*, 36(4), pp. 311-317.
- VOSS, U. *et al.* (2014): “Induction of self awareness in dreams through frontal low current stimulation of gamma activity”, *Nature Neuroscience*, 17 (6), pp. 810-812.